

The University of Chicago

CANONICAL EXPRESSIONS IN BOOLEAN ALGEBRA

A DISSERTATION SUBMITTED TO THE
FACULTY OF THE DIVISION OF THE
PHYSICAL SCIENCES IN CANDIDACY FOR
THE DEGREE OF DOCTOR OF PHILOSOPHY

DEPARTMENT OF MATHEMATICS

1937

By

ARCHIE BLAKE

Private Edition, Distributed by
THE UNIVERSITY OF CHICAGO LIBRARIES
CHICAGO, ILLINOIS

1938

The University of Chicago

CANONICAL EXPRESSIONS IN BOOLEAN ALGEBRA

A DISSERTATION SUBMITTED TO THE
FACULTY OF THE DIVISION OF THE
PHYSICAL SCIENCES IN CANDIDACY FOR
THE DEGREE OF DOCTOR OF PHILOSOPHY

DEPARTMENT OF MATHEMATICS

1937

By
ARCHIE BLAKE

Private Edition, Distributed by
THE UNIVERSITY OF CHICAGO LIBRARIES
CHICAGO, ILLINOIS

1938



TABLE OF CONTENTS

	Page
INTRODUCTION	1
CHAPTER I <u>Preliminaries</u>	
Section	
1. Rigor in mathematics	3
2. Expressions	5
3. Notation	6
4. Definition of Boolean algebra	7
CHAPTER II <u>Finite Boolean Expressions</u>	
5. Introduction	10
6. Preliminary definitions and reduction	11
7. Monomials	16
8. Polynomials--elementary properties and development . .	18
9. Formal inclusion and the syllogistic and simplified forms	24
10. Properties of the syllogistic form	29
11. Example	37
CHAPTER III <u>Propositional Functions</u>	
12. Introduction	45
13. Notations and process of transformation	46
14. Existence of an evaluation in which all reduced transforms take the value $\textcircled{1}$	53
15. Existence of an evaluation reducing the initial transform to $\textcircled{1}$	55
16. The equivalence theorem	57
17. Example	58



Digitized by the Internet Archive
in 2026

https://archive.org/details/IA41548403_0043

CANONICAL EXPRESSIONS IN BOOLEAN ALGEBRA

INTRODUCTION

It is the purpose of this paper to study the transformation of expressions by means of the laws of Boolean algebra--that is, the algebra of logic--and to examine the properties of certain canonical forms to which expressions may be reduced by means of these transformations.

The first chapter explains the notations and the assumptions about matters of logic and the foundations of mathematics that are used in this paper. The point of view is better described as intuitive-mechanistic than postulational. The principle merit of the postulational method lies in the mechanization of the theory to which it is applied; it does not leave the underlying logic in any more satisfactory state than before. In this paper the interest is in mechanization for its own sake rather than merely as an attribute of the postulational method; the point of view is intuitive as regards the logical foundations.

A postulational definition of Boolean algebra is given in Chapter I. This treatment is inadequate to the needs of the transformation theory to be presented in this paper. For this purpose we desire the intuitive notion of an expression. This concept and the notations used in dealing with it are explained in the first chapter.

Chapter II deals with the transformation of finite Boolean expressions, that is, expressions not involving infinite

sums or products, by means of the Boolean laws. Much of this theory is well known, and is stated here briefly without proof. The sylogistic and simplified forms here given are believed to be new; the use of them provides essentially a combination of the method of Boole and Schröder with that of Poretsky. The use of the simplified canonical form is illustrated in Section 11 on a classical problem of Boole.

When the class of expressions is extended to represent propositional functions by means of infinite sums and products, many new and as yet unsolved problems arise. It is the purpose of Chapter III to define and develop some of the properties of a process of transformation applicable in this calculus. This method, though of general applicability, is here studied only for the simplest case in which quantification is confined to the individuals (that is, to the Prädikatenkalkul or engere Funktion-enkalkul).

For any given expression the process defines an infinite sequence of expressions called the transforms of the given expression; associated with each of these transforms is a Boolean expression in the sense of Chapter II, the reduced transform. The sequence of reduced transforms serves as a canonicalization of the given expression in the following sense: if there exists a reduced transform which is equivalent to the universal, the given expression is reduced to $\bigvee$ in any evaluation; otherwise, on any infinite class there exists an evaluation for which the given expression takes the value $\bigwedge$. The method yields as corollaries Löwenheim's theorem that a formula of the Prädikatenkalkul is fulfillable on a denumerable class if at all and Skolem's completeness theorem, the latter being obtained without use of the Skolem reduction.

An illustrative example is worked out in Section 17.

CHAPTER I

PRELIMINARIES

1. Rigor in mathematics. If a sharp distinction is made between the postulates of a mathematical theory on the one hand and the logical axioms, rules of inference, or other principles which we use to guide our thought on the other, it is clear that, as regards the rigor of the theory obtained, the postulational method as commonly used serves merely to transfer the philosophical difficulties from one part of the theory to another. Mathematically the main utility of the method lies in the mechanization which it brings about. Mechanization sought for its own sake, however, need perhaps not be confined to the postulational method. For example, it has been characteristic of that method to proceed from denumerable sets of undefined terms and postulates by discrete steps to denumerable sets of theorems. It is not clear that mathematics must remain within these bounds; for example a probability theory might be possible in which the advance of inference from nescience to theorem would be continuous rather than discrete.

Two points of view are used in this paper. A postulational definition of Boolean algebra is given in Section 4, but this definition is not adequate for the theory to be obtained. For the complete development of the transformation theories of Chapters II and III it appears necessary to use the notion of an expression; this concept is explained in the next section.

The study of logic by symbolic methods presents one complication which is absent from most postulational theories: a principle of logic may play a dual rôle as part of the system being studied and one of the rules of inference used in the study. Three procedures are possible here, according as it is permitted to use the principle for inference throughout the theory, only after its formal equivalent has been established (if ever), or not at all. The second and third of these procedures are subject to the difficulty of our separating off in thought a part of logic from the rest of it with sufficient definiteness to be sure that we never make tacit use of the principle in question. The intuitive first procedure, while it does not throw so much light on the logical foundations, is free from the difficulty just mentioned and is, of course, the most powerful for the purpose of developing the superstructure of the subject; it is the method of this paper.

We assume the whole of ordinary logic, including the law of the excluded middle, the existence of finite and infinite classes of objects for thought, the principle that every infinite class contains a denumerable subclass, the principle of choice, the validity of mathematical induction, the properties of equality or identity, and the notions of a function, of a construction (which is a set of directions for replacing one object for thought by another), and of an expression, which is to be explained in the next section.

It is sought to avoid circular fallacies by a method which was much used by E. H. Moore, namely to restrict the arguments of functions to fixed classes which are well-defined in terms of given classes, rather than merely to the complete domains of apparent significance.

Certain of the foregoing logical principles, notably those that are concerned with infinite classes, could perhaps be dispensed with for the purposes of Chapter II, and there are parts of the argument of Chapter III which appear capable of formulation in such a way as to avoid the use of the principle of choice. But this is not attempted in the present paper, for reasons already mentioned.

2. Expressions. The term "formula," as it occurs in the writings of workers in symbolic logic, appears often to be intended to denote a material object--some combination of symbols, written, printed, or spoken. To avoid confusion with this concept (which is more restrictive than is needed in the present theory) we use the term expression to denote an object for thought such as is intended to be conveyed by a formula. It is a formula apart from its fortuitous peculiarities of size, shape, location, and identity, and its existence as a material object.

Thus $a + b$ and $b + a$ considered as formulas are not the same, because they are different material objects, but considered as expressions they are equal--that is, the same. On the other hand the expressions $a + b$ and $b + a$ are not equal. In fact, an expression such as $a + b$ is a set of instructions explaining how a number is to be obtained from two numbers, a and b , and a function, $+$; the expression exists independently of any particular numbers or function. Since many functions, $+$, could be defined which are not symmetric, it is clear that $a + b$ and $b + a$, considered as expressions, are not equal, though they would in many number-systems represent the same number; in such cases it is appropriate to call the expressions equivalent in some sense. Various kinds of equivalence of expressions will be used in this paper.

Different kinds of expressions are useful in different contexts. Descriptions will be given at the beginnings of the next two chapters of the kinds of expressions to be discussed and the notations adopted for them.

3. Notation. The theorems will be stated in words or in symbols or both, according to convenience. The symbolism used is that of E. H. Moore,¹ modified, to suit the needs of the present theory, in the following respects: The symbol $\sim$, used by Moore for logical equivalence, is here used for a defined equivalence of expressions, logical equivalence being denoted by $\equiv$, punctuated appropriately with dots. The latter symbol is used in Moore's work for equivalence by definition; definitions will here be indicated by use of the subscript D applied to $=$ or to $\equiv$. The capital German letters are reserved for classes and the small Greek letters for functions as in Moore's notation, but the special rôles which he assigned to certain letters are replaced by others, which will be introduced as needed. The convention that all given classes are non-null is not used in this paper.

Most of Moore's notations will be familiar to workers in symbolic logic; there are a few that need brief explanation: The symbol $\ni$ is used to mean such that and $|$ for uniquely. The possession of a property is indicated by writing the symbol for the property as a superscript applied to the symbol for the object which possesses the property; thus a^P means a has the property P, and the assertion for every x, x has the property P is

¹Moore's notation, part of which he obtained from Peano and part developed in his own work, is described in General Analysis, Part I, by Eliakim Hastings Moore with the cooperation of Raymond Walter Barnard, Memoirs of the American Philosophical Society, vol. 1 (1935), pp. 28-38.

symbolized by $x \cdot$ or x^P or by $x^P(x)$. The implication symbol is used to express both quantification and implication. Variables for which no quantification is indicated are understood to be free variables, whose scope is the whole assertion.

The notation $x \mathcal{A}$ means x is a member of the class $\mathcal{A}$.

The symbols I, II, III, and IV applied to classes denote respectively the properties of having exactly one member, having only a finite number of members, being denumerably infinite, and being non-denumerably infinite.

To represent a function ϕ defined on a class $\mathcal{P}$ and with functional values all belonging to a class $\mathcal{Q}$ Moore used the notation ϕ on $\mathcal{P}$ to $\mathcal{Q}$ or $(\phi(p)|p \in \mathcal{P})$. The class of the values of this function is symbolized by use of square brackets instead of round: $[\phi(p)|p \in \mathcal{P}]$.

The functions dealt with in this paper are understood to be single-valued unless the contrary is indicated. Assertions and definitions in which multiple-valued functions enter are understood to hold for any determination (the same throughout the assertion) of the value of the function; in such cases the statement, if written in words, will refer to "a" value of the function rather than "the" value.

4. Definition of Boolean algebra. The meaning of the term "Boolean algebra"¹ is rather definitely agreed upon by writers on symbolic logic, except for certain matters of formulation

¹A good introduction to this subject is given by Louis Couturat in L'Algebra de la Logique, 1905. An English translation of this book by Lydia Gillingham Robinson was published in 1914 by the Open Court Publishing Company. A bibliography of symbolic logic for the period from 1666 to 1935 inclusive is given by Alonzo Church in the Journal of Symbolic Logic, vol. 1 (1936), p. 121.

such as whether elements of the algebra are to be spoken of as equal or merely equivalent, and whether $+$, $\times$, etc., are to be regarded as functions or operators of some other kind.

We use two points of view. First, in this section, Boolean algebra is defined by means of functions and equality (which is used in the sense of identity--about the same as the concept of *13.01 in Principia Mathematica). In Chapter II we deal with expressions instead of functions; expressions which can be transformed into each other according to the Boolean rules will be called equivalent, but they will not be called equal unless they are the same expression in the sense explained in Section 2.

We define a Boolean algebra as a class $\mathcal{B} = [b]$, two special members of the class, $\mathbb{V}$, called universal, and $\mathbb{0}$, called null, two functions, $\oplus$ and $\otimes$, each on $\mathcal{B}\mathcal{B}$ to $\mathcal{B}$, and another function, $\circ$, on $\mathcal{B}$ to $\mathcal{B}$, provided they satisfy certain conditions:

(4.1) Definition.

$$(\mathcal{B} = [b], \mathbb{V}^{\mathcal{B}}, \mathbb{0}^{\mathcal{B}}, \oplus \text{ on } \mathcal{B}\mathcal{B} \text{ to } \mathcal{B}, \otimes \text{ on } \mathcal{B}\mathcal{B} \text{ to } \mathcal{B}, \circ \text{ on } \mathcal{B} \text{ to } \mathcal{B})_{\text{Bool alg}} \equiv_{\mathcal{D}}.$$

$(\mathcal{B}, \mathbb{V}, \mathbb{0}, \oplus, \otimes, \circ)$ properties (1) ... (7) below

$$(1) \quad \mathbb{V} \oplus \mathbb{V}^{\circ} = \mathbb{V}$$

$$(2) \quad \mathbb{0} = \mathbb{V}^{\circ}$$

$$(3) \quad (b_1 \otimes b_2) \oplus (b_1 \otimes b_2^{\circ}) = b_1$$

$$(4) \quad b_1 \oplus b_2 = b_2 \oplus b_1$$

$$(5) \quad b_1 \otimes b_2 = (b_1^{\circ} \oplus b_2^{\circ})^{\circ}$$

$$(6) \quad (b_1 \oplus b_2) \oplus b_3 = b_1 \oplus (b_2 \oplus b_3)$$

$$(7) \quad \mathbb{V}^{\circ} \neq \mathbb{V}$$

In this definition the letter b connotes "member of $\mathcal{B}$ " as indicated in the notation " $\mathcal{B} = [b]$." The German capital $\mathcal{B}$ will be reserved throughout this paper for Boolean algebras, and the small Roman b will be reserved except in Section 11 for an element of a Boolean algebra.

The use of the two elements $\mathbb{V}$ and $\mathbb{A}$ and the three functions $\oplus$, $\otimes$, and $\odot$ as primitives is dictated by convenience in connection with Chapter II. Properties (1), ..., (7) are obtained by adapting Huntington's fourth set of postulates¹ to the present choice of primitives.

Huntington's postulate 4.5 is redundant, as he has proved.² His other postulates, together with properties which characterize $\mathbb{V}$, $\mathbb{A}$, and $\otimes$, which are not among his primitives, are embodied in the present definition. Boolean algebra in the present sense is an instance of Huntington's, but the converse is not true because equality as here used is more restrictive than in the sense of Huntington's postulates A, B, C, and D.

The relation of inclusion between elements of a Boolean algebra is here defined in terms of equivalence:

4.2) Definition.

$$b_1 \otimes b_2 \equiv_D b_1 \odot \oplus b_2 = \mathbb{V}$$

¹E. V. Huntington, New sets of independent postulates for the algebra of logic, with special reference to Whitehead and Russell's Principia Mathematica, Transactions of the American Mathematical Society, vol. 35 (1933), p. 274.

²E. V. Huntington, Boolean algebra. A correction, Transactions of the American Mathematical Society, vol. 35 (1933), p. 557.

CHAPTER II

FINITE BOOLEAN EXPRESSIONS

5. Introduction. This chapter deals with Boolean algebra considered as a theory of the transformation of expressions. This point of view has been used in the past; it is here formulated somewhat more elaborately than usual. Sections 6, 7, and 8 are occupied with the elementary theory of the transformation of expressions according to the Boolean rules; a large part of this theory is well known. The theory of formal inclusion and the syllogistic form presented in Sections 9 and 10 is believed to be new.

In this chapter we are given a class, $\mathcal{U}$, of letters, $[a]$, and five symbols, $\vee$, $\wedge$, $+$, $\times$, and $'$, called respectively universal, null, plus, times, and negative; the last three are also called function-symbols. The theory deals with the expressions corresponding to the finite formulas--i.e. formulas in which the functions appear only a finite number of times--obtained by treating $'$ as a unary and $+$ and $\times$ as binary functions on $[\mathcal{U}, \vee, \wedge]$ to the same class. The class of all such expressions will be denoted by $\mathfrak{E} = [e]$. The cases where no function-symbols appear are included, so that a , $\vee$, and $\wedge$ are expressions, as well as $a_1 + a_2$, $(a_1 \times a_2)'$, etc. Such a combination of symbols as $a_1 \times a_2 \times a_3$ is not an expression, because the manner of association of the factors is not indicated; the expressions $(a_1 \times a_2) \times a_3$ and $a_1 \times (a_2 \times a_3)$ are not the same. Most of the theorems to be obtained hold for any manner of

associating the factors in a product or the terms in a sum. In such cases the sums and products will be written without parentheses, the convention being that parentheses are to be supplied in any significant manner. In cases where no confusion can arise the symbol $\times$ will be omitted, the notation $a_1 a_2$ being written instead of $a_1 \times a_2$.

It is understood that all letters appearing in the statement of a theorem, whether explicitly or not, are members of $\mathcal{U}$. The notation $\mathcal{U}_0^{\text{incl } e}$ will be used to mean that the class of letters $\mathcal{U}_0$ includes all letters appearing in e , and $\mathcal{U}_0^{\text{min incl } e}$ to mean that $\mathcal{U}_0$ is the minimal class with this property.

For every defined term a dual term also is understood to be defined. The dual of a definition is obtained from it by interchanging $\vee$ with $\wedge$ and $+$ with $\times$, and replacing every previously defined term which appears in it by its dual. The dual of each defined symbol or term will be written by use of an asterisk applied to it as a superscript. The dual theory will be understood without mention except in a few cases of special interest.

6. Preliminary definitions and reduction. The letters and their negatives will be called co-letters, and the class of them denoted by $\mathcal{C} = [c]$. By the conjugate of a co-letter is meant its negative if it is a letter, otherwise, the letter whose negative it is. The conjugate of a co-letter c will be written c^+ , and conjugate co-letters will be said to have opposite signs. (It should be noted that while the negative and the conjugate of a letter are the same, the negative of a' is a'' , which is not the same as $a'^+ = a$. There are three important differences between $'$ and $+$: $'$ is a given function-symbol applicable to any expression;

$\dagger$ is a defined function on the class of co-letters.)

Two co-letters will be called conjunctive if they are either equal or conjugates; the symbol $=^\dagger$ will be used for this relation.

These definitions are symbolized as follows:

(6.1) Definition.

$$(1) \quad \mathcal{C} =_D [c] =_D [\text{all co-letters}] =_D [\text{all } a, a']$$

$$(2) \quad c^\dagger =_D \begin{pmatrix} a' & | & c = a \\ a & | & c = a' \end{pmatrix}$$

$$(3) \quad c_1 =^\dagger c_2 \equiv_D: c_1 = c_2 \cdot \vee \cdot c_1 = c_2^\dagger$$

Two expressions will be called equivalent in case they can be transformed into each other (in a finite number of steps) by means of the Boolean rules. The rules of transformation used will be those of Huntington's fourth set, which was used in Section 4, except that a stronger property than his postulate 4.0 is used; Huntington's 4.0 will be obtained in Theorem (6.72) as a corollary. In this formulation, in contrast to that of Section 4, we need Huntington's properties A, B, C, and D of equality; property C is here used in extended form in connection with our strengthening of Huntington's 4.0.

For convenience in defining equivalence, which we symbolize by $\sim$, it is desirable first to define three preliminary relations, $\sim_1$, $\sim_2$, and $\sim_3$:

(6.2) Definition.

$$(1) \quad e_1 \sim_1 e_2 :\equiv_D:$$

$$e_1 = e_2$$

$$\cdot^v. e_1 = V + V' \cdot e_2 = V$$

$$\cdot^v. e_1 = \Lambda \cdot e_2 = V'$$

$$\cdot^v. \exists e \ni e_1 = e_2 e + e_2 e'$$

$$\cdot^v. \exists (e_3, e_4) \ni e_1 = e_3 + e_4 \cdot e_2 = e_4 + e_3$$

$$\cdot^v. \exists (e_3, e_4) \ni e_1 = e_3 e_4 \cdot e_2 = (e_3' + e_4')'$$

$$\cdot^v. \exists (e_3, e_4, e_5) \ni e_1 = (e_3 + e_4) + e_5$$

$$\cdot e_2 = e_3 + (e_4 + e_5)$$

$$(2) \quad e_1 \sim_2 e_2 :\equiv_D: e_1 \sim_1 e_2 \cdot^v. e_2 \sim_1 e_1$$

$$(3) \quad e_1 \sim_3 e_2 :\equiv_D: \exists (e_3^{\text{part of } e_1} \cdot e_4^{\text{part of } e_2}) \ni$$

$$e_3 \sim_2 e_4 \cdot e_2 \text{ is obtained from } e_1 \text{ by}$$

$$\text{replacing } e_3 \text{ by } e_4$$

$$(4) \quad e_1 \sim e_2 :\equiv_D: e_1 \sim_3 e_2 \cdot^v.$$

$$\exists (n^{\text{integer}} \geq 3, e_3, \dots, e_n) \ni e_1 \sim_3 e_3$$

$$\cdot e_j \sim_3 e_{j+1} (j = 3, \dots, n-1) \cdot e_n \sim_3 e_2$$

The essential difference between this formulation and Huntington's is that, with the exception of his postulate 4.0, which provides that the system have at least two inequivalent elements, all of his postulates are statements of sufficient conditions for equivalence, whereas the equivalence of the present definition is characterized by the necessary and sufficient conditions given.

We define inclusion in terms of equivalence:

 (6.3) Definition.

$$e_1 < e_2 :\equiv_D: e_1' + e_2 \sim V$$

The following theorem is useful in many applications:

(6.4) Theorem. Equivalence of expressions is preserved under the process of identifying letters or of replacing any letter throughout by $\vee$ or $\wedge$.

The fact that equality is preserved under this process is part of our intuitive idea of an expression. Then by the first part of Definition (6.2) equivalence in the first preliminary sense is preserved, and by the other parts of the definition equivalence itself is preserved.

The following duality theorem expresses a well-known property of Boolean algebra:

(6.5) Theorem.

$$e_1 \sim e_2 \cdot \equiv \cdot e_1 \sim^* e_2$$

For the purpose of studying the relation of the present transformation theory to Boolean algebra in the sense of Section 4 we define, for any function ξ on $\mathcal{A}$ to any Boolean algebra, its extension to the class $\mathcal{E}$ of all expressions, in the following manner:

(6.6) Definition.

ξ on $\mathcal{A}$ to $\mathcal{B}$.). $\xi(e) =_D$ the member of $\mathcal{B}$ obtained from e by replacing $\vee$, $\wedge$, $+$, $\times$, and $'$ respectively by $\bigvee$, $\bigwedge$, $\oplus$, $\otimes$, and $\odot$, and each letter, a , by $\xi(a)$.

On comparing the definitions of (6.2) and (4.1) it is found that the equivalences among expressions are transformed by the extension of any function on $\mathcal{A}$ to $\mathcal{B}$ into equalities among the members of $\mathcal{B}$:

(6.7) Theorem.

$$\xi \text{ on } \mathcal{A} \text{ to } \mathcal{B} \cdot e_1 \sim e_2 \cdot) \cdot \xi(e_1) = \xi(e_2)$$

A converse will be obtained in Theorem (8.2).

(6.71) Corollary.

$$\xi \text{ on } \mathcal{A} \text{ to } \mathcal{B} \cdot e_1 < e_2 \cdot) \cdot \xi(e_1) \otimes \xi(e_2)$$

(6.72) Corollary.

$$\vee \approx \wedge$$

For if $\vee \sim \wedge$, then in any Boolean algebra $\mathbb{V} = \mathbb{A}$. This corollary is the statement of Huntington's postulate 4.0 which was to be obtained from our formulation (6.2).

In the proofs of the ensuing theorems the parts which can be supplied by use of well-known Boolean rules derivable (as in Huntington's paper) from the definition of equivalence will usually be omitted. Some of the converse proofs will be given in more detail.

The object of this paper is to transform given expressions into special forms which are suitable for studying certain properties of the expressions. The transformations are to be made in such a way as to preserve equivalence. We first reduce the given expression to a form in which the function-symbol ', if present, applies only to letters. This is accomplished by means of the rules $e'' \sim e$, $(e_1 + e_2)' \sim e_1'e_2'$, and $(e_1e_2)' \sim e_1' + e_2'$. When this has been done, the resulting expression is to be expanded by means of the distributive laws $e_1(e_2 + e_3) \sim e_1e_2 + e_1e_3$ and $(e_1 + e_2)e_3 \sim e_1e_3 + e_2e_3$ and thus reduced to polynomial form. Before discussing this form it will be desirable to develop certain properties of monomials.

7. Monomials. By a monomial will be meant an expression in which the function-symbol ', if present, applies only to letters and the function-symbol + is not present. The class of monomials will be denoted by $\mathcal{M} = [m]$, and the property of being a monomial by M . It is not in general possible to reduce a given expression to monomial form.

A co-letter c , or $\vee$ or $\wedge$, will be said to be a factor of a monomial m if it appears as an argument to a $\times$ -symbol in m or if it is the whole expression m . This relation will be symbolized by the notation $c^{\text{fact } m}$ or $\vee^{\text{fact } m}$ or $\wedge^{\text{fact } m}$.

The following five theorems on the relations between monomials and their factors are useful in subsequent proofs.

(7.1) Theorem.

$$m \sim \wedge \quad \equiv \quad \wedge^{\text{fact } m} \cdot \vee \cdot \exists c \ni c^{\text{fact } m} \cdot c^+ \text{ fact } m$$

The fact that either of the conditions on the right is sufficient to ensure that $m \sim \wedge$ is obvious. To prove the converse, suppose neither of these conditions satisfied, and consider any Boolean algebra, $\mathcal{B}$. We can choose ξ on $\mathcal{U}$ to $\mathcal{B}$ so that for every co-letter factor c of m , $\xi(c) = \mathbb{O}$, whence $\xi(m) = \mathbb{O}$, contrary to $m \sim \wedge$ and Theorem (6.7).

(7.2) Theorem.

$$m < c \quad \equiv \quad m \sim \wedge \cdot \vee \cdot c^{\text{fact } m}$$

The sufficiency of the condition on the right is obvious. Conversely, if $m < c$, $mc^+ \sim \wedge$, whence by Theorem (7.1) either $m \sim \wedge$ or $\exists$ some co-letter and its conjugate which are factors of mc^+ but not of m .

(7.21) Corollary.

$$mc^{\dagger} \sim \wedge \cdot c^{-\text{fact } m} \cdot m \sim \wedge$$

The next theorem is a generalization of Theorem (7.2).

(7.3) Theorem.

$$(1) \quad m_2 \not\sim \wedge : c^{\text{fact } m_2} \cdot c^{\text{fact } m_1} : ; m_1 < m_2$$

$$(2) \quad m_1 \not\sim \wedge \cdot m_1 < m_2 : ; c^{\text{fact } m_2} \cdot c^{\text{fact } m_1}$$

Part (1) is obvious, and part (2) follows from $m_1 < m_2 < c$ by Theorem (7.2).

(7.4) Theorem.

$$m_1 c < m_2 : \equiv : m_1 c < m_2 c : :$$

$$c^{\dagger} \text{ fact } m_1 \cdot m_2 < c \cdot m_1 < m_2$$

The equivalence is obvious. To prove the implication, suppose $c^{\dagger} - \text{fact } m_1 \cdot m_2 \nless c \cdot m_1 \not\sim \wedge$. Then $m_1 c \not\sim \wedge$, and, by part (2) of Theorem (7.3), every co-letter factor of m_2 is a factor of $m_1 c$; such a co-letter must be c or else a factor of m_1 . But c is not a factor of m_2 because $m_2 \nless c$. Then, since $m_2 \not\sim \wedge$, $m_1 < m_2$ by part (1) of Theorem (7.3).

(7.5) Theorem.

$$m < m_1 + m_2 \cdot c_1 \neq c_2^{\dagger} \cdot c_1^{-\text{fact } m} \cdot c_2^{-\text{fact } m} \\ \cdot m_1 < c_1 \cdot m_2 < c_2 \cdot m \sim \wedge$$

For $mc_1^{\dagger}c_2^{\dagger} < m_1c_1^{\dagger}c_2^{\dagger} + m_2c_1^{\dagger}c_2^{\dagger}$, and both terms on the right are equivalent to $\wedge$. Then if $\wedge^{-\text{fact } m}$, some co-letter c and its conjugate must be factors of $mc_1^{\dagger}c_2^{\dagger}$ by Theorem (7.1). Also $c \neq^{\dagger} c_1$, for if so $c_1^{\text{fact } mc_2^{\dagger}}$, contrary to hypothesis; similarly $c \neq^{\dagger} c_2$. Therefore c and $c^{\dagger}$ are factors of m , and consequently $m \sim \wedge$.

8. Polynomials--elementary properties and development.

By a polynomial will be meant a monomial or a finite sum of monomials. The class of polynomials will be denoted by $\mathcal{P} = [p]$, and the property of being a polynomial by P . A monomial will be said to be a term of a polynomial in case it is the whole polynomial or an addend in the polynomial. The notation $m^{\text{term } p}$ will be used to mean that m is a term of p .

Every expression can be reduced to an equivalent polynomial.

(8.1) Theorem.

$$\exists p \sim e$$

The proof was indicated at the end of Section 6. (The method there used for constructing the polynomial could be made unequivocal, if we wished, by specifying the order in which the various processes are to be performed, so that a given expression would give rise to a unique polynomial.)

The reduction to polynomial form enables us to prove a converse of Theorem (6.7); by combining these two results we obtain the equivalence theorem, showing the relation between the present theory of the transformation of expressions and Boolean algebra as defined in Section 4: The equivalences among expressions are the identical formulas in the sense that they transform into equalities for any Boolean algebra $\mathcal{B}$ and any function on $\mathcal{A}$ to $\mathcal{B}$.

(8.2) Theorem.

$$\mathcal{B} : e_1 \sim e_2 \equiv: \xi^{\text{on } \mathcal{A} \text{ to } \mathcal{B}} .). \quad \xi(e_1) = \xi(e_2)$$

The implication to the right was proved in Theorem (6.7). Next, suppose $e_1 \not\sim e_2$; then $e_1 e_2' + e_1' e_2 \not\sim \Lambda$. Consider any $p \sim e_1 e_2' + e_1' e_2$; then some term, m , of p is $\not\sim \Lambda$. Choose $\xi(c) = \mathbb{V}$ for every c fact m ; then by Theorem (7.1) $\xi(m) = \mathbb{V}$, whence $\xi(p) = \mathbb{V}$ by Theorem (6.71), and $\xi(e_1 e_2' + e_1' e_2) = \mathbb{V}$ by Theorem (6.7). Then since $\xi(e_1 e_2' + e_1' e_2) = \xi(e_1) \otimes (\xi(e_2))^\odot \oplus (\xi(e_1))^\odot \otimes \xi(e_2)$, we have $\xi(e_1) \neq \xi(e_2)$.

(8.21) Corollary.

$\mathcal{B} :: e_1 < e_2 :: \xi$ on $\mathcal{A}$ to $\mathcal{B}$.). $\xi(e_1) \otimes \xi(e_2)$

There is an infinitude of polynomials equivalent to any given expression; thus $a_1 + a_2$ and $a_2 + a_1$ are equivalent, as are also Λ and $\Lambda + \Lambda$, as well as a_1 and $a_1 + a_1 a_2$, etc. Certain of the ways in which equivalent expressions can differ, namely those that depend on the reordering and reassociating of addends and multipliers, are not important for the purposes of this paper;¹ as a device for ignoring these differences we make the following definition:

Two expressions will be called congruent ($\cong$) in case one of them can be transformed into the other by means of the commutative and associative laws of addition and multiplication alone. In the definitions of canonical expressions to be given we shall content ourselves with properties sufficient to characterize the expressions in the sense of congruence. (These properties could be strengthened sufficiently to characterize a unique expression by defining an order relation on the class of letters and adopting arbitrary conventions about ordering and association.)

¹But there exist purposes for which they are important.

Evidently congruence implies equivalence, is self-dual, and is reflexive, symmetrical, and transitive. Two monomials are congruent if and only if they have the same factors each repeated the same number of times. The relation of congruence divides the class of monomials into a set of mutually exclusive and exhaustive subclasses, the maximal classes of mutually congruent monomials. Two polynomials are congruent in case the number of terms from each of these classes in one of them is the same as in the other.

The process of absorption, by which irrelevant factors are dropped from monomials and terms from polynomials, is next to be described.

Each monomial which is equivalent to $\vee$ or $\wedge$ is to be replaced by $\vee$ or $\wedge$ as the case may be. From the other monomials $\wedge$ will be absent, and all occurrences of $\vee$ and all repetitions of letters are to be eliminated by the rules $\vee\vee \sim \vee$, $c\vee \sim c \sim \vee c$, and $cc \sim c$. A monomial such as is thus obtained will be said to be absorptive (abs):

(8.3) Definition.

$$\begin{aligned} m^{\text{abs}} &:=_{\text{D}}: m \sim \wedge \cdot m = \wedge \\ &: \vee^{\text{fact } m} \cdot m = \vee \\ &: \text{no co-letter appears twice as a factor of } m \end{aligned}$$

By $\text{abs}(m)$ will be meant a monomial obtained from m by the construction indicated above. The function abs is multiple-valued because no choice of the order and manner of association of the factors has been made, but this multiplicity is not important for the present purposes because the values of $\text{abs}(m)$ are mutually congruent. In fact any two absorptive monomials which

are equivalent must also be congruent, for if they are not both $\wedge$ they have the same co-letter factors by part (2) of Theorem (7.3).

(8.5) Theorem.

$$m_1^{\text{abs}} \sim m_2^{\text{abs}} \Rightarrow m_1 \equiv m_2$$

To extend the definition of absorption to a polynomial, p , denote by $\mathcal{M}_0$ the class of its monomial terms and let the $n \geq 1$ mutually exclusive and exhaustive non-null subclasses into which $\mathcal{M}_0$ is divided by the relation of equivalence be denoted by $\mathcal{M}_j$ ($j = 1, \dots, n$). Delete each of the $\mathcal{M}_j$ whose members are included in the members of another of the $\mathcal{M}_j$, and delete from each of the remaining $\mathcal{M}_j$ all but one of its members. Let the remaining monomials be $m_1, m_2, \dots, m_k$, and form a sum of all of $\text{abs}(m_1), \text{abs}(m_2), \dots$, and $\text{abs}(m_k)$, obtaining a value of $\text{abs}(p)$, as required.

(8.6) Definition.

$$\begin{aligned} p^{\text{abs}} &:=_{\text{D}}; \text{mterm } p \Rightarrow m^{\text{abs}} \\ &: (m_1 \text{ term } p \cdot m_2 \text{ term } p) \text{ distinct as terms of } p \\ &\Rightarrow m_1 \not\equiv m_2 \end{aligned}$$

(The notation $(m_1 \cdot m_2)$ distinct as terms of p means that m_1 and m_2 occur at different places in the expression p , not that $m_1 \neq m_2$.)

It should be noted that the set of laws of transformation which are included in the absorption process is not self-dual. In fact, the replacing of terms which are equivalent to $\wedge$ by $\vee$ often involves use of the law¹ $cc^+ \sim \wedge$. The dual of this law

¹The use of this law constitutes an extension of the present application of the term "absorption" beyond its usual meaning.

plays a special rôle in the present formulation, which will be brought out in the ensuing discussion.

It is evident from the definition of absorption and Theorem (8.5) that the values of $\text{abs}(p)$ are mutually congruent. The analogue of Theorem (8.5) itself, that two absorptive polynomials which are equivalent must be congruent, is not true, as we see by the example of $a + a'$ and $\vee$.

This theorem is true, however, if one member of the equivalence is $\wedge$: The only absorptive polynomial which is equivalent to $\wedge$ is $\wedge$ itself. This fact affords a simple and much-used method of verifying any equivalence or inclusion whose truth-value is to be found, for any equivalence or inclusion can be reduced to the form $e \sim \wedge$ and therefore to $p \sim \wedge$ and to $\text{abs}(p) \sim \wedge$.

More convenient at times when the problem is given in the form of an inclusion to be verified, say $e_1 < e_2$, is the following procedure: We reduce e_1 to polynomial form and e_2 to polynomial* form, say to p and p^* respectively, and subject p to the absorption construction and p^* to its dual, so that the original inclusion takes the form

$$\text{abs}(p) < \text{abs}^*(p^*).$$

Since a sum is included in a product if and only if every addend in the sum is included in every factor in the product, it suffices to compare each monomial term of $\text{abs}(p)$ with each monomial* factor of $\text{abs}^*(p^*)$. These comparisons can be made by inspection, for a monomial $m \not\sim \wedge$ is included in a monomial* $m^* \not\sim \vee$ if and only if m has a co-letter factor which is an addend in m^* . This method is self-dual.

An important classical method of examining an inclusion $e_1 < e_2$ is the method of dichotomy or development. A polynomial $p \not\sim \wedge$ is said to be developed with respect to a letter a in case, for every term m of p , either a or a' is a factor of m , but not both. A monomial m may be developed with respect to a letter a by means of the relation $m \sim ma + ma'$. To verify an inclusion $e_1 < e_2$ we replace it by $p_1 < p_2^{abs}$, and develop p_1 with respect to all the letters that appear in p_2 . A term m_1 of this developed expression is included in a term $m_2 \not\sim \wedge$ of p_2 in case no letter enters with opposite signs in m_1 and m_2 . These and other properties of the developed form are well known.

The problem of determining, for two given expressions, e_1 and e_2 , whether $e_1 < e_2$, is generalized, in the work of Boole and Schroder, to the problem of solving Boolean equations, together with the concomitant problems of elimination. More general results are obtained in the method of Poretsky,¹ who seeks all the consequences and causes of a set of data, rather than only those obtainable by choosing some set of variables to be eliminated and some set to be solved for. While his results are quite complete, his method is impracticably long in all but the simplest problems. In fact, in a problem whose data involve n letters, the number of consequences to be read from a Poretskian table may in extreme cases be as high as $2^{2^n(2^n-1)} \sim 2^{2^n-1}$.

The object of the next section is to develop the properties of a sylogistic form to which a given expression can be reduced. This form enables us to reach essentially the results

¹Platon Poretsky, Sept lois fondamentales de la théorie des égalités logiques, Bulletin de la Society Physico-Mathématique de Kasan, ser. 2, vol. 8 (1898), pp. 33-103, 129-181, and 183-216. See especially pages 142 and 150 for a description of the table of consequences.

of Poretsky much more efficiently than by use of his tables of consequences and causes.

9. Formal inclusion and the syllogistic and simplified forms. A polynomial p_1 will be said to be formally¹ included in a polynomial p_2 in case every term of p_1 is included in some term of p_2 . This relation will be written $p_1 \ll p_2$. Two polynomials which are formally included in each other will be said to be formally equivalent ($\approx$):

(9.1) Definition.

$$(1) \quad p_1 \ll p_2 : \equiv_D \text{ } m_1^{\text{term } p_1} \text{ .) . } \exists m_2^{\text{term } p_2} \text{ } \exists . m_1 < m_2$$

$$(2) \quad p_1 \approx p_2 : \equiv_D \text{ } p_1 \ll p_2 \cdot p_2 \ll p_1$$

(9.11) Corollary.

$$p_1 \ll p_2 \cdot \text{) } \cdot p_1 < p_2$$

In case a polynomial p_1 is included in another polynomial p_2 , but not formally included in it, p_1 will be said to be in-formally included in p_2 ($p_1 \leq p_2$).

(9.2) Definition.

$$p_1 \leq p_2 : \equiv_D \text{ } p_1 < p_2 \cdot p_1 \ll p_2$$

For example, $V \leq a + a'$.

A polynomial will be said to be syllogistic (syl) in case it admits no informally included polynomial, that is, in case every polynomial which is included in it is also formally included in it:

¹This term has been used in different senses in the past. It was used by Peirce in connection with the distinction between "material" and "formal" implication, and more recently by Whitehead and Russell for the inclusion of propositional functions.

(9.3) Definition.

$$p^{\text{syl}} \equiv_D p_1 \nmid p(p_1)$$

 (9.31) Corollary.

$$p^{\text{syl}} \equiv: p_1 < p \cdot) \cdot p_1 \ll p$$

The reason for the choice of the term "syllogistic" in this connection will be seen by an example: The syllogistic result to be deduced from $b_1 \otimes b_2$ and $b_2 \otimes b_3$ is $b_1 \otimes b_3$. If the premises are written as equations, $b_1 b_2^{(1)} = \mathbb{A}$ and $b_2 b_3^{(1)} = \mathbb{A}$, they may be combined into the single equation $b_1 b_2^{(1)} \oplus b_2 b_3^{(1)} = \mathbb{A}$. The left member here admits the informally included term $b_1 b_3^{(1)}$. If we adjoin this term, we obtain $b_1 b_2^{(1)} \oplus b_2 b_3^{(1)} \oplus b_1 b_3^{(1)} = \mathbb{A}$; the expression on the left now admits no informally included term. Since this property was reached by means of syllogistic inference applied to the given data (and this holds in general, as we shall see), it is appropriate to speak of the resulting expression as being syllogistic.

The importance of the syllogistic form is that it provides a simple characterization of the class of included expressions. The investigation whether one polynomial is formally included in another can be performed at sight (by means of Theorem (7.3)), for it involves a mere comparison between the monomial terms of the two polynomials. Therefore if a given expression be transformed into an equivalent polynomial and this polynomial reduced to syllogistic form, the formally included polynomials (which are the same as all included polynomials) can be read off directly, so that all the Poretskian consequences of the given data are embodied in usable form in the syllogistic polynomial.

This method depends on the fact that there exists a syllogistic polynomial equivalent to any given expression, e . For consider any $\mathcal{U}_0^{\text{incl } e}$ and let $\mathcal{M}_0$ denote the class of absorptive monomials included in e and involving only letters of $\mathcal{U}_0$, $\mathcal{M}_0 =_{\mathcal{D}} [\text{all } m^{\text{abs}} < e \ni \mathcal{U}_0^{\text{incl } m}]$. Choosing one monomial from each of the mutually exclusive and exhaustive non-null subclasses into which $\mathcal{M}_0$ is divided by the relation of equivalence, and forming a sum of all these monomials, we obtain a polynomial which is equivalent to e and evidently has the syl property. Such a polynomial will be said to be in the complete canonical form with respect to $\mathcal{U}_0$ ($\text{com}_{\mathcal{U}_0}$), and, in case $\mathcal{U}_0^{\text{min incl } e}$, will be written $\text{com}(e)$. The function com is multiple-valued, but the values are mutually congruent.

(9.4) Definition.

$$\begin{aligned} \mathcal{U}_0^{\text{incl } p} &:: p^{\text{com}_{\mathcal{U}_0}} \equiv_{\mathcal{D}}: \\ &\quad m^{\text{term } p} \cdot m^{\text{abs}} \\ &:: m_1 < p \ni \mathcal{U}_0^{\text{incl } m_1} \cdot \exists m_2^{\text{term } p} \sim m_1 \end{aligned}$$

(9.41) Corollary.

- (1) $\mathcal{U}_0^{\text{incl } e} \cdot (\text{com}(e))^{\text{com}_{\mathcal{U}_0}} \cdot P$
- (2) $\mathcal{U}_0^{\text{incl } p} \cdot p^{\text{com}_{\mathcal{U}_0}} \cdot p^{\text{syl}}$
- (3) $\exists p^{\text{syl}} \sim e$; for example, $\text{com}(e)$
- (4) $\mathcal{U}_0^{\text{incl } (p_1, p_2)} \cdot p_1^{\text{com}_{\mathcal{U}_0}} \sim p_2^{\text{com}_{\mathcal{U}_0}} \cdot p_1 \equiv p_2$

The foregoing method of constructing the complete canonical form for a given expression is far from being the most efficient means of obtaining a syllogistic expression equivalent to it. The best form for this purpose is the simplified canonical form, defined as follows:

(9.5) Definition.

$$p^{\text{sim}} \cdot \Xi_D \cdot p^{\text{syl} \cdot \text{abs}}$$

To prove the existence of such an expression it suffices to demonstrate that the syl property is retained in the process of absorption:

(9.6) Theorem.

$$(1) \quad p_1 \ll p_2 + p_3 \cdot p_3 \ll p_2 \cdot \cdot \cdot p_1 \ll p_2$$

$$(2) \quad (p_1 + p_2)^{\text{syl}} \cdot p_2 \ll p_1 \cdot \cdot \cdot p_1^{\text{syl}}$$

$$(3) \quad p^{\text{syl}} \cdot \cdot \cdot (\text{abs}(p))^{\text{syl}}$$

To prove (1) consider any term, m , of p_1 , and suppose $m \nless p_2$. Then $\exists m_1^{\text{term}} p_3 \ni m < m_1$, and since $p_3 \ll p_2$, $\exists m_2^{\text{term}} p_2 \ni m_1 < m_2$, whence $m < m_2$, $m \ll p_2$. The second part of the theorem is proved similarly, and the third part follows from the second together with the definition of absorption, for the terms deleted in absorption are all formally included in the sum of terms retained.

(9.7) Theorem.

$$\exists p^{\text{sim}} \sim e$$

In fact if we start with any $p_1^{\text{syl}} \sim e$, and subject it to the absorption construction, we obtain the desired p^{sim} .

We shall use the notation $\text{sim}(e)$ to denote any one of the values of $\text{abs}(\text{com}(e))$.

The next theorem shows in what way the simplified form is best for the purpose for which the syllogistic form is intended: the simplified form is minimal within any class of equivalent syllogistic polynomials.

(9.8) Theorem.

$$p_1^{\text{sim}} \sim p_2^{\text{syl}} ; m_1^{\text{term}} p_1 \cdot \cdot \cdot \exists m_2^{\text{term}} p_2 \sim m_1$$

Since $m_1 < p_1 < p_2$, and p_2^{syl} , $\exists m_2^{\text{term}} p_2 \ni m_1 < m_2$. Similarly $\exists m_3^{\text{term}} p_1 \ni m_2 < m_3$, whence $m_1 < m_3$ and therefore, by Definition (8.6), $m_1 = m_3$, whence $m_1 \sim m_2$.

(9.9) Theorem.

$$p_1^{\text{sim}} \sim p_2^{\text{sim}} \cdot \cdot \cdot p_1 \dot{=} p_2$$

For each of the maximal sets of mutually equivalent monomials, p_1 and p_2 have the same number, zero or one, of members of the set among their terms, by Definition (8.6) and Theorem (9.8). The respective terms are congruent by Theorem (8.5), and therefore $p_1 \dot{=} p_2$.

On comparing Definition (9.4) with Theorem (9.8) it is seen that the complete and simplified forms are opposites in the sense that the complete form is maximal syllogistic while the simplified form is minimal syllogistic. The syllogistic form itself has an opposite in a sense, namely the developed form. The relation between these two, and a number of properties of syllogistic expressions, including several methods of constructing the syllogistic form, are discussed in the next section. It suffices to construct the syllogistic instead of the simplified form, for the latter can then be reached by absorption.

10. Properties of the syllogistic form. It is evident that all monomials are syllogistic, and all binomials of which one term is equivalent to $\vee$ or to $\wedge$. For a binomial $m_1 + m_2$ in which neither m_1 nor m_2 is equivalent to $\wedge$ we define two classes of co-letters: $\mathcal{C}_1$ is the class of all co-letters, c ,

such that c appears as a factor in m_1 or m_2 (or both) while $c^\dagger$ does not appear as a factor in either m_1 or m_2 . $\mathfrak{C}_2$ is the class of all co-letters, c , such that c appears as a factor of m_1 and $c^\dagger$ appears as a factor of m_2 .

In terms of these classes we define two functions, σ and γ , of the binomial. We define also a relation to be called "seminclusion" and symbolized by $\sqsubset$: if a monomial is not equivalent to Λ , a polynomial is semincluded in it provided every co-letter factor of the monomial is a factor of some term of the polynomial. If the polynomial is itself a monomial which is not equivalent to Λ , the relation reduces to ordinary inclusion.

(10.1) Definition.

- (1) $m_1 \not\sim \Lambda \not\sim m_2$:):
 $\mathfrak{C}_1^{\text{null}} \cdot \sigma(m_1 + m_2) =_{\mathcal{D}} \vee$
 $\mathfrak{C}_1^{-\text{null}} \cdot \sigma(m_1 + m_2) =_{\mathcal{D}}$ a monomial whose factors
 are the members of $\mathfrak{C}_1$ each used once
- (2) $m_1 \not\sim \Lambda \not\sim m_2 \cdot \mathfrak{C}_2^I$:).
 $\gamma(m_1 + m_2) =_{\mathcal{D}}$ the $c \ni c^{\text{fact } m_1} \cdot c^\dagger \text{ fact } m_2$
- (3) $m \not\sim \Lambda$!):
 $p \sqsubset m \equiv_{\mathcal{D}}: c^{\text{fact } m} \cdot \supset m_1^{\text{term } p} \ni c^{\text{fact } m_1}$

(10.11) Corollary.

- $m_1 \not\sim \Lambda \not\sim m_2$:):
- (1) The values of $\sigma(m_1 + m_2)$ are mutually congruent.
- (2) $\sigma(m_1 + m_2) \not\sim \Lambda \cdot m_1 + m_2 \sqsubset \sigma(m_1 + m_2)$
- (3) $\sigma(m_1 + m_2) < m \cdot m \not\sim \Lambda \cdot m_1 + m_2 \sqsubset m$
- (4) $\mathfrak{C}_2^I \cdot m_1 \not\sim \vee \not\sim m_2 \not\sim m_1$

The notation " σ " is used to suggest the word "syllogism," for the term $\sigma(m_1 + m_2)$ is the syllogistic result to be obtained from the terms m_1 and m_2 . The next theorem states this and the fact that $\sigma(m_1 + m_2)$ is the maximal such result; it also provides in terms of the co-letter factors of m_1 and m_2 a necessary and sufficient condition that $(m_1 + m_2)^{\text{syl}}$.

(10.2) Theorem.

- $$\begin{aligned}
 & (1) \quad (m_1 + m_2)^{-\text{syl}} \\
 \equiv: & (2) \quad m_1 \not\wedge \not\wedge m_2 \cdot \mathfrak{C}_2^I \\
 \equiv: & (3) \quad m_1 \not\wedge \not\wedge m_2 \cdot \sigma(m_1 + m_2) \leq m_1 + m_2 \\
 :): & (4) \quad m_1 \not\wedge \not\wedge m_2 \cdot m < \sigma(m_1 + m_2) \quad (m \leq m_1 + m_2)
 \end{aligned}$$

It will suffice to prove $(2) \cdot) \cdot (3) \cdot) \cdot (1) \cdot) \cdot (2) \cdot (4)$.

To prove (3) from (2) we note first that $\sigma(m_1 + m_2) < m_1 + m_2$ by the law $e_1 e_2 < e_1 c + e_2 c^\dagger$, and second, that $\sigma(m_1 + m_2)$ is not included in either m_1 or m_2 by part (2) of Theorem (7.3). Condition (1) follows from (3) by Definition (9.3).

To prove conditions (2) and (4) from (1) consider any $m \leq m_1 + m_2$. Evidently $m \not\wedge$ and $m_1 \not\wedge \not\wedge m_2$. Define $\mathfrak{C}_3 =_{\text{D}} [\text{all } c^{\text{fact } m_1} \text{--fact } m]$ and $\mathfrak{C}_4 =_{\text{D}} [\text{all } c^{\text{fact } m_2} \text{--fact } m]$. Since $m \not\wedge m_1 \not\wedge$, $\mathfrak{C}_3^{-\text{null}}$ by part (1) of Theorem (7.3), and similarly $\mathfrak{C}_4^{-\text{null}}$. Then by Theorem (7.5), $c_3^{\mathfrak{C}_3} \cdot c_4^{\mathfrak{C}_4} \cdot)$. $c_3 = c_4^\dagger$, whence $\mathfrak{C}_3^I$, $\mathfrak{C}_4^I$, and $\mathfrak{C}_3 \subset \mathfrak{C}_2$. To prove $\mathfrak{C}_2 \subset \mathfrak{C}_3$, suppose $\exists c^{\mathfrak{C}_2} \text{--}\mathfrak{C}_3$, so that $c^{\text{fact } m}$ and $m < c$. Then $m < m_1 c + m_2 c$, and since $m_2 c \sim \wedge$, $m < m_1 c < m_1$, contrary to $m \leq m_1 + m_2$. All members of $\mathfrak{C}_1$ and their conjugates are absent from $\mathfrak{C}_2$ and therefore also from $\mathfrak{C}_3$ and $\mathfrak{C}_4$. Therefore all co-letter factors of $\sigma(m_1 + m_2)$ are absent from $\mathfrak{C}_3$ and $\mathfrak{C}_4$, and, since they are all factors of m_1 or m_2 , they are factors of m by the definitions

of $\mathfrak{G}_3$ and $\mathfrak{G}_4$, so that $m < \sigma(m_1 + m_2)$.

With regard to the relation of $\sigma(m_1 + m_2)$ to m_1 and m_2 , it is clear from the foregoing theorem that if no letter enters with opposite signs in m_1 and m_2 , then $\sigma(m_1 + m_2)$ is included in either term, while if more letters than one have this property, then $\sigma(m_1 + m_2) \not\leq m_1 + m_2$. If there is just one such letter, $\sigma(m_1 + m_2)$ is the maximal informally included monomial.

The next theorem is of fundamental importance. If a polynomial is not syllogistic, it contains a binomial which is not syllogistic and whose σ -function is informally included in the polynomial.

(10.3) Theorem.

$$p^{-\text{syl}} \cdot \exists (m_1, m_2)^{\text{terms } p} \ni (m_1 + m_2)^{-\text{syl}} \cdot \sigma(m_1 + m_2) \leq p$$

Let n be the number of distinct letters appearing in p . Define $\mathcal{M}_0 =_D [\text{all } m^{\text{abs.} \leq p}]$. Define the degree of any member of $\mathcal{M}_0$ as the number of its distinct co-letter factors. Let m be any member of $\mathcal{M}_0$ of maximal degree; this degree is less than n because a developed expression is formally included in any expression which includes it. There is therefore some letter, a , appearing in p and absent from m , and the monomials ma and ma' , being of higher degree than m , are both $\ll p$, say $ma < m_1^{\text{term } p}$. $ma' < m_2^{\text{term } p}$. Since $m \leq m_1 + m_2$, $(m_1 + m_2)^{-\text{syl}}$; also since $m < \sigma(m_1 + m_2)$ by Theorem (10.2), we have $\sigma(m_1 + m_2) \not\leq p$, for otherwise $m \ll p$, contrary to the way m was chosen; therefore $\sigma(m_1 + m_2) \leq p$.

(10.31) Corollary.

$$m_1 \not\wedge \not\wedge m_2 \cdot (m_1 + m_2 + \sigma(m_1 + m_2))^{\text{syl}}$$

This theorem, together with Theorems (9.6) and (10.2), and part (3) of Corollary (9.41), provides a simple means of reducing a given expression to syllogistic form. For if a polynomial is not syllogistic, then two of its terms can be found whose σ -function may be adjoined to the original polynomial and augments the class of inequivalent terms that are formally included; the class of all such terms to be reached is finite because such terms are all terms of the complete canonical form. Finally, by part (1) of Theorem (9.6), the work may be simplified at any stage by absorption without decreasing the class of formally included terms. This is one of the most efficient ways of reducing an expression to simplified form. It is desirable to eliminate terms by absorption wherever possible, and to combine the simpler terms before the more complex. An example will be given in Section 11.

The developed and syllogistic forms may be contrasted in two ways. First, a developed expression is formally included in every polynomial in which it is included, while a syllogistic expression includes formally every polynomial which it includes. Second, the process of reducing a given polynomial to developed form consists in replacing monomials, m , by binomials $ma + ma'$, while the construction of the syllogistic form proceeds by the adjunction of new terms $\sigma(m_1 + m_2)$; the two constructions involve the equivalence $e_1 \sim e_1e_2 + e_1e_2'$, using it in opposite senses.

The syl property, like the property of development, can be defined for individual co-letters:

(10.4) Definition.

- (1) $p^{\text{syl}_c} :=_{\mathcal{D}}:$
 $(m_1 \cdot m_2)^{\text{terms } p} \ni (m_1 + m_2)^{-\text{syl}} \cdot \gamma(m_1 + m_2) =^+ c \cdot).$
 $\sigma(m_1 + m_2) \ll p$
- (2) $p^{\text{syl}_{\mathcal{C}_0}} :=_{\mathcal{D}} p^{\text{syl}_c} (c \mathcal{C}_0)$

(10.41) Corollary.

$$\mathcal{A}_0 \text{ incl } p :): p^{\text{syl}_{\mathcal{A}_0}} =. p^{\text{syl}}$$

Evidently a polynomial having the property syl_c can be constructed which is equivalent to any given polynomial, for any p^{syl} has this property by the foregoing corollary, but it is not necessary to obtain the property syl in order to reach syl_c . In fact, the desired result can be reached simply by adjoining to the given polynomial all σ -terms obtainable from pairs of terms of the given polynomial by eliminating c :

(10.5) Theorem. By the adjunction of all informally included σ -terms obtainable from pairs of terms of a polynomial p by eliminating c , p is reduced to an equivalent polynomial having the property syl_c .

In fact, neither c nor c^+ is a factor of any of the new terms adjoined, while the σ -function obtained by eliminating c from any pair of the old terms is formally included in the new polynomial.

The next theorem is analogous to Theorem (9.6): the property syl_c is preserved under absorption.

(10.6) Theorem.

- (1) $(p_1 + p_2)^{\text{syl}_c} \cdot p_2 \ll p_1 \cdot). p_1^{\text{syl}_c}$
 (2) $p^{\text{syl}_c} \cdot). (\text{abs}(p))^{\text{syl}_c}$

Part (1) is proved by reference to part (1) of Definition (10.4). For p_1 provides a smaller class of pairs (m_1, m_2) than $p_1 + p_2$, while every monomial which is $\ll p_1 + p_2$ is also $\ll p_1$ by part (1) of Theorem (9.6), so that $p_1^{\text{syl}c}$ follows from $(p_1 + p_2)^{\text{syl}c}$. The second part of the theorem follows from the first by the definition of absorption.

The next theorem is useful in proving the important Theorem (10.8).

(10.7) Theorem.

$p^{\text{syl}c} : (m_1, m_2)^{\text{terms } p} \ni m_1 \nless c \cdot m_2 \nless c^\dagger$
 · no letter other than c or c[†] enters with opposite signs
in m₁ and m₂ :
 $\exists m^{\text{term } p} \ni m \nless c \cdot m \nless c^\dagger \cdot m_1 + m_2 \sqsubset m$

If $m_1 \nless c^\dagger$ or $m_2 \nless c$, then m_1 or m_2 , as the case may be, serves as the desired m . Otherwise $(m_1 + m_2)^{-\text{syl}c}$, and, since $p^{\text{syl}c}$, $\exists m^{\text{term } p} \ni \sigma(m_1 + m_2) < m$, so that c and $c^\dagger$ are absent from m and every co-letter factor of m is a factor of $\sigma(m_1 + m_2)$ and therefore of m_1 or m_2 .

It will next be proved that an expression which is syl-
logistic with respect to a co-letter remains so when a σ -term
obtained by eliminating any co-letter from a pair of its terms
is adjoined.

(10.8) Theorem.

$p^{\text{syl}c_1} \cdot (m_1, m_2)^{\text{terms } p} \ni (m_1 + m_2)^{-\text{syl}c_2} \cdot$
 $(p + \sigma(m_1 + m_2))^{\text{syl}c_1}$

If $c_1 =^\dagger c_2$, the conclusion is obvious because $\sigma(m_1 + m_2)$ is then free of c_1 and $c_1^\dagger$. Take $c_1 \neq^\dagger c_2$ and denote $\sigma(m_1 + m_2)$

by m_3 . Supposing the conclusion of the theorem to be false we have

$$\exists m_4 \text{ term } p \ni (m_3 + m_4)^{-\text{syl}c_1} \cdot \sigma(m_3 + m_4) \nless p.$$

One of m_1 and m_2 has c_2 as a factor--let it be m_1 without loss of generality; then $c_2^+ \text{ fact } m_2$. Similarly let $c_1^+ \text{ fact } m_3$ and $c_1^+ \text{ fact } m_4$.

Since $m_3 + c_2 \subseteq m_1$ and $c_1^+ \text{ -fact } m_3$, it follows that $c_1^+ \text{ -fact } m_1$, and similarly $c_1^+ \text{ -fact } m_2$. Also $c_1^+ \text{ -fact } m_4$, and not both of c_2 and c_2^+ are factors of m_4 .

If $c_2^+ \text{ -fact } m_4$, Theorem (10.7) applied to m_1 and m_4 yields $\exists m_5 \text{ term } p \ni c_1^+ \text{ -fact } m_5 \cdot c_1^+ \text{ -fact } m_5 \cdot m_1 + m_4 \subseteq m_5$; also $c_2^+ \text{ -fact } m_5$. If also $c_2^+ \text{ fact } m_4$, it follows from part (1) of Definition (10.1) that $c_2^+ \text{ fact } \sigma(m_3 + m_4)$ and $\sigma(m_3 + m_4) < m_5$. But if $c_2 \text{ -fact } m_4$, then by Theorem (10.7) applied to m_2 and m_4 $\exists m_6 \text{ term } p \ni c_1^+ \text{ -fact } m_6 \cdot c_1^+ \text{ -fact } m_6 \cdot m_2 + m_4 \subseteq m_6$; also $c_2 \text{ -fact } m_6$. Then by Theorem (10.7) applied to m_5 and m_6 $\exists m_7 \text{ term } p \ni m_7 \nless c_2 \cdot m_7 \nless c_2^+ \cdot m_5 + m_6 \subseteq m_7$. Since neither c_1 nor c_1^+ is a factor of m_5 or m_6 , neither of them is a factor of m_7 , and therefore, by part (1) of Definition (10.1), every co-letter factor of m_7 is a factor of $\sigma(m_3 + m_4)$, so that $\sigma(m_3 + m_4) << p$.

This theorem and Theorem (10.6) provide a useful means of constructing a syllogistic form for a given polynomial. The eliminations may be performed letter by letter, and when once they have been completed for a letter, a , so that the expression is syl_a , this property is retained throughout the rest of the construction.

Another theorem which is often useful in constructing a syllogistic form for a given expression is the following:

(10.9) Theorem. If two polynomials have the property syl_c , any polynomial obtained by expanding their product by means of the distributive laws $e_1(e_2 + e_3) \sim e_1e_2 + e_1e_3$ and $(e_1 + e_2)e_3 \sim e_1e_3 + e_2e_3$ has also the property syl_c .

Denote the given polynomials and the expansion of their product by p_1 , p_2 , and p respectively. Suppose the conclusion false; then $\exists (m_1, m_2) \text{ terms } p \ni (m_1 + m_2)^{-\text{syl}_c} \cdot \sigma(m_1 + m_2) < p$. Let c fact m_1 , $c^\dagger$ fact m_2 . Also m_1 was obtained as the product of a term m_{11} of p_1 and a term m_{12} of p_2 . Neither of these is $\sim \wedge$ nor has $c^\dagger$ as a factor. Similarly

$$\exists (m_{21}^{\text{term } p_1} \nless c, m_{22}^{\text{term } p_2} \nless c).$$

Since $\sigma(m_1 + m_2) + c \sqsubset m_1 < m_{11}$, we have $\sigma(m_1 + m_2) + c \sqsubset m_{11}$, and similarly $\sigma(m_1 + m_2) + c^\dagger \sqsubset m_{21}$. Then no letter other than c or $c^\dagger$ enters with opposite signs in m_{11} and m_{21} , for if so it would follow that $\sigma(m_1 + m_2) \sim \wedge$. By Theorem (10.7),

$\exists m_3^{\text{term } p_1} \ni m_3 \nless c \cdot m_3 \nless c^\dagger \cdot m_{11} + m_{21} \sqsubset m_3$; therefore

$\sigma(m_1 + m_2) < m_3$, by part (1) of Theorem (7.3). Similarly

$\exists m_4^{\text{term } p_2} \ni \sigma(m_1 + m_2) < m_4$. Then since one of the terms of p is equivalent to m_3m_4 we have $\sigma(m_1 + m_2) << p$.

Two corollaries should be mentioned. First, the theorem holds if the property syl_c is replaced throughout by syl , by Corollary (10.41). Second, if an expression has the properties p^* and abs^* , any polynomial obtained by expanding it will have the property syl , since every monomial* with the property abs^* has the property syl . This method was known to Peirce and

his students.¹

The next theorem is useful in applications of the theory:

(10.95) Theorem. If a polynomial p_1 has the property syl_c , any polynomial, p_2 , obtained from p_1 by dropping all terms in which a or a' appears as a factor has also the property syl_c .

For if $p_2^{-\text{syl}_c}$, $\exists (m_1, m_2)^{\text{terms } p_2} \ni (m_1 + m_2)^{-\text{syl}_c}$
 $\cdot \sigma(m_1 + m_2) \nless p_2$. Since m_1 and m_2 are free of a , $\sigma(m_1 + m_2)$ is so also, and therefore any monomial in which $\sigma(m_1 + m_2)$ is included is free of a . Now m_1 and m_2 are terms of p_1 so that $\sigma(m_1 + m_2) \ll p_1$, and the term of p_1 in which $\sigma(m_1 + m_2)$ is included is free of a , by part (2) of Theorem (7.3), and is therefore a term of p_2 , whence $\sigma(m_1 + m_2) \ll p_2$.

The theorem evidently holds also if the property syl_c is replaced throughout by syl .

11. Example. Principles of logic sufficient to solve any problem in finite Boolean algebra have been known at least since the time of Aristotle. The contribution of the modern logicians has been the use of algebraic methods which have made possible the surmounting of obstacles of complexity which alone prevented the ancients from reaching all the modern results. The merit of a method in finite Boolean algebra lies, therefore, not primarily in the scope of the problems which it will solve, but in the simplicity with which it solves the most important problems in the subject.

The methods of Boole and Schröder, inspired by the example of ordinary algebra, are admirably suited to the obtaining

¹It is mentioned at several places in Studies in Logic by members of the Johns Hopkins University, 1883.

of certain results, namely the resultants of elimination and the solutions of equations. When further results are desired, the methods of solution available have not been so satisfactory. The method of Poretsky reaches all the consequences of a given inclusion, but inefficiently. The methods of the present paper effect some improvement in this regard, and are here illustrated on a classical problem of Boole.

The problem chosen is Example 5 in Chapter IX of Boole's Laws of Thought.¹ This problem is discussed at length by Schröder,² who calls it his Prüfstein für die Methoden.³

Certain conventions of notation that have been used in developing the foregoing theory are of very little use in applying it, and are here abandoned: The letters a, b, c, and e will no longer be reserved for special uses, but will be used for elements of a Boolean algebra in the sense of Section 4, and also (since any class of objects for thought may serve this purpose) as letters to be used in forming expressions. The symbols $\vee$, $\wedge$, $+$, $\times$, and $'$ will serve a similar double purpose. No confusion will be possible because the copulas = and $\oplus$ will be used in this section exclusively for elements of a Boolean algebra.

The problem of Boole is then formulated as follows:
Given the equation

$$(11.1) \quad \begin{aligned} & a'c'(bd + b'd' + e') + ade'(bc' + b'c) \\ & + a(b + e)(cd + c'd') + (a' + b'e')(cd' + c'd) = \Lambda. \end{aligned}$$

¹George Boole, An Investigation of the Laws of Thought, 1854, pp. 146-149. (This problem appears on pp. 155-158 in vol. 2 of the 1916 edition of Boole's Collected Logical Works published by the Open Court Publishing Company.)

²Ernst Schröder, Vorlesungen über die Algebra der Logik, vol. 1, 1890, pp. 522-528, 562-569, 571, 586-589, and 591.

³Ibid., p. 568.

It is required to determine what independent relations among b , c , and d may be inferred from this equation, and what relations among these three quantities may be inferred with the further hypothesis $a = V$; similarly what relations among a , c , and d may be inferred from the equation, and what with the further hypothesis $b = V$.

Since none of the conclusions to be obtained concerns the symbol e , the first step in the solutions of Boole and Schröder is the elimination of this quantity. Next, to determine whether any independent relations exist among b , c , and d , the quantity a is eliminated, and to find what conclusions follow from $a = V$ (as well as the converse question), the procedure is to solve for a . The questions asked with regard to a , c , and d are answered in the same way.

Schröder's solution leaves little to be desired for simplicity and directness in answering the questions asked, but it does not provide a convenient characterization of the class of all conclusions which can be inferred from equation (11.1). Poretsky's table of consequences¹ is intended for this purpose. The developed form for the expression (11.1) has 21 terms, and the Poretsky table of consequences has therefore 2^{21} rows and 2^{11} columns. The number of equations to be read off from each column (excluding equations of the form $x = x$, and not counting such equations as $x = y$ and $y = x$ as distinct) is $\frac{1}{2} \cdot 2^{21}(2^{21} - 1)$, and since the number of columns is 2^{11} , the total number of consequences is $2^{52} - 2^{31}$.

¹Platon Poretsky, Sept lois fondamentales de la théorie des égalités logiques, Bulletin de la Society Physico-Mathématique de Kasan, ser. 2, vol. 8 (1898), pp. 33-103, 129-181, and 183-216. See especially pages 142 and 150 for a description of the table of consequences.

The method of this paper accomplishes the same result more efficiently. We transform the expression in equation (11.1) into an equivalent expression in simplified form; this expression, while it does not exhibit all the Poretskian conclusions explicitly, enables us to read them off by immediate (as distinguished from syllogistic) inference.

Expanding the left member of equation (11.1) we obtain the polynomial

$$\begin{aligned}
 & a'bc'd + a'b'c'd' + a'c'e' + abc'de' + ab'cde' \\
 (11.2) \quad & + abcd + abc'd' + acde + ac'd'e + a'cd' + a'c'd \\
 & + b'cd'e' + b'c'de'.
 \end{aligned}$$

It is desirable at each stage of the work to delete all terms that are absorbed by other terms of the expression. Accordingly, before starting to adjoin new terms, we drop the term $a'bc'd$, which is absorbed by $a'c'd$. It is also very helpful to combine the pairs of simple terms before the more complex. Accordingly we arrange the terms in order of complexity:

$$\begin{aligned}
 & a'c'e' + a'cd' + a'c'd + a'b'c'd' + abcd + abc'd' \\
 (11.3) \quad & + acde + ac'd'e + b'cd'e' + b'c'de' + abc'de' + ab'cde'.
 \end{aligned}$$

The process of reducing this expression to simplified form proceeds now by the adjunction of σ -terms which are informally included in the expression thus far obtained, and the dropping of terms which are absorbed.

The first two terms may be combined, in accordance with Theorem (10.2). They give rise to the new term

$$\sigma(a'c'e' + a'cd') \doteq a'd'e',$$

and since this term is not included in any term already obtained,

it is adjoined. No term of the expression is included in this new term, and there is therefore no absorption to be performed. No further combination of pairs of three-letter terms thus far obtained is possible; thus the terms $a'c'e'$ and $a'c'd$ can not be combined because no letter enters them with opposite signs, while the terms $a'cd'$ and $a'c'd$ can not because two letters enter them with opposite signs. (See part (2) of Theorem (10.2).) We proceed to the combination of three- and four-letter terms.

The terms $a'c'e'$ and $abc'd'$ give rise to the term $bc'd'e'$, which we adjoin. The terms $a'c'e'$ and $b'cd'e'$ give rise to $a'b'd'e'$, but this is included in the term $a'd'e'$ already obtained, and is therefore not adjoined. The terms $a'cd'$ and $a'b'c'd'$ give rise to $a'b'd'$, which we adjoin; this term absorbs the term $a'b'c'd'$, which is therefore to be deleted.

Before proceeding with the combining of three- and four-letter terms it is well to compare the new three-letter term with those that are already obtained. It is found that the terms $a'c'd$ and $a'b'd'$ give rise to the term $a'b'c'$; no further new combinations of three-letter terms are possible.

The terms $a'b'd'$ and $ac'd'e$ give rise to $b'c'd'e$, which we adjoin. All combinations of the three- and four-letter terms thus far obtained have been completed, and we proceed to combine the four-letter terms among themselves. It is found that no pair of the four-letter terms can be combined to give anything new. We proceed to the three- and five-letter terms.

The terms $a'c'e'$ and $abc'd'e'$ give rise to the term $bc'd'e'$, which absorbs $abc'd'e'$. The terms $b'c'd'e'$ and $bc'd'e'$ give rise to the term $c'd'e'$, which absorbs them both. The terms $c'd'e'$ and $abcd$ give rise to $abde'$. The terms $c'd'e'$ and $abc'd'$

give rise to $abc'e'$. The terms $c'de'$ and $bc'd'e'$ give rise to the term $bc'e'$, which absorbs $bc'd'e'$ and $abc'e'$.

The terms $c'de'$ and $ab'cde'$ give rise to $ab'de'$, which absorbs the term $ab'cde'$. The terms $abcd$ and $ab'de'$ give rise to the term $acde'$, which combines with $c'de'$ to give the term ade' , which absorbs $acde'$, $abde'$, and $ab'de'$. The terms ade' and $acde$ give rise to acd , which absorbs the terms $abcd$ and $acde$. The terms ade' and $b'cd'e'$ give rise to $ab'ce'$.

No further σ -terms can be adjoined, and no further absorptions are possible; the expression is now reduced to the simplified form

$$\begin{aligned}
 (11.4) \quad & a'c'e' + a'cd' + a'c'd + a'd'e' + a'b'd' \\
 & + a'b'c' + c'de' + bc'e' + ade' + acd \\
 & + abc'd' + ac'd'e + b'cd'e' + b'c'd'e + ab'ce'.
 \end{aligned}$$

The questions asked about the original equation (11.1) can now be answered by inspection. There are no independent relations among b , c , and d , because every term of the expression (11.4) involves either a or e . The expression does contain terms free of b and e , however, and the conclusion about a , c , and d independently is

$$(11.41) \quad a'cd' + a'c'd + acd = \wedge,$$

which agrees with the results of Boole and Schröder.

Next, what can be inferred from the further datum $a = \vee$? This is equivalent to $a' = \wedge$. If this inquiry had been the only one in the problem, it would have been desirable to simplify the work as much as possible by using the datum at the beginning of the calculation. As it is we adjoin the term a' to the expression (11.4). This term absorbs all other terms having a' as a

factor and serves to eliminate a from the rest, giving the expression

$$(11.42) \quad \begin{aligned} & a' + c'de' + bc'e' + de' + cd + bc'd' + c'd'e \\ & + b'cd'e' + b'c'd'e + b'ce', \end{aligned}$$

which must be syllogistic since (11.4) was. The term $c'de'$ is absorbed by de' , the term $b'cd'e'$ by $b'ce'$, and the term $b'c'd'e$ by $c'd'e$. The expression becomes

$$(11.43) \quad a' + de' + cd + bc'e' + bc'd' + c'd'e + b'ce'.$$

Since it is desired to obtain results independent of e , we restrict attention to the terms in which e does not appear, obtaining the result

$$(11.44) \quad a' + cd + bc'd' = \wedge.$$

Similarly, to investigate the consequences of $b = \vee$, we adjoin the term b' to (11.4), and, after simplifications similar to those performed above on (11.41), we obtain

$$(11.45) \quad \begin{aligned} & b' + c'e' + a'cd' + a'c'd \\ & + a'd'e' + ade' + acd + ac'd'. \end{aligned}$$

From this we obtain the desired result by ignoring the terms involving e :

$$(11.46) \quad b' + a'cd' + a'c'd + acd + ac'd' = \wedge.$$

Schröder's conclusion from $b = \vee$ is $a' + c + d = \vee$, i.e., $ac'd' = \wedge$. This is the last term in equation (11.46); the other terms are in Schröder's point of view not regarded as important.

The solutions of Boole and Schröder obtain lower as well as upper bounds for a and b . These also can be found simply from

the expression (11.4). To solve $x \subseteq a$ we add a to the expression (11.4) as a' was added before. We obtain

$$(11.5) \quad a + c'e' + cd' + c'd + d'e' + b'd' + b'c',$$

and if it is desired to ignore the terms involving e this reduces to

$$(11.51) \quad a + cd' + c'd + b'd' + b'c'.$$

The original equation is thus reduced to

$$(11.52) \quad a = a + cd' + c'd + b'd' + b'c',$$

and the solutions of $x \subseteq a$ are to be read off directly from this; they are the Boolean entities corresponding to the expressions which are formally included in the right member.

It is seen that the results obtainable from the expression (11.4) are much more than the special part of them that results when attention is restricted to the terms that are free of e . Had we been content to eliminate e at the beginning of the calculation, the work of reaching the simplified form for the remaining expression would have been considerably reduced.

It should be remarked that there are many formulas which can be studied more efficiently than by reduction to syllogistic form. For example, the simplified form for the equations

$$(11.6) \quad a = b = c$$

is

$$(11.7) \quad bc' + b'c + ac' + a'c + ab' + a'b = \wedge,$$

and the use of this equation would in many applications be less conducive to simplicity than the device of replacing b and c throughout by a .

CHAPTER III

PROPOSITIONAL FUNCTIONS

12. Introduction. Many new problems arise when the class of Boolean formulas is augmented by the admittance of infinite sums and products. The purpose of this chapter is to define and develop some of the properties of a process of transformation, analogous to a process which is much used by mathematicians in constructing proofs, which is applicable to many problems in the theory of propositional functions. Attention is here confined to expressions of the first order, that is of the erste Stufe of the Prädikatenkalkül in the terminology of Hilbert and Bernays¹ or of the engere Funktionenkalkül in the terminology of Hilbert and Ackermann²--that is, to formulas in which the only variables that are subject to quantification are the individuals.

For any expression of this calculus the process generates an infinite sequence of transforms which are equivalent to the given expression in the sense of always yielding the same result on evaluation to a 2-element Boolean algebra. Associated with each transform is a finite Boolean expression, the reduced transform; the sequence of reduced transforms serves as a canonicalization of the given expression in the following sense:

¹D. Hilbert and P. Bernays, Grundlagen der Mathematik, 1934.

²D. Hilbert and W. Ackermann, Grundzüge der Theoretischen Logik, 1928. This and the book just mentioned serve as a good introduction to this part of the subject.

If any reduced transform is equivalent to $\neg V$, the given expression takes the value $\bigvee$ in any evaluation. Otherwise there exists an evaluation in which all the reduced transforms take the value $\bigwedge$, and in this case there exists an evaluation on any infinite class which reduces the initial transform and therefore the given expression to $\bigwedge$.

The method does not provide a general Entscheidungs-verfahren because no bound is obtained for the number of reduced transforms that may have to be examined before the first one is reached which is equivalent to V , nor general procedure of any kind for proving that none of them can be equivalent to V . (The method of Gödel¹ has been extended by Church² to the result that the general case of the Entscheidungsproblem of the engere Funktionenkalkül is unsolvable; the precise meaning of these results, however, is still somewhat a matter of controversy.³)

13. Notations and process of transformation. The expressions of this chapter are certain combinations of the following given symbols and classes of symbols:

- (1) An infinite class $\mathcal{X} = [x]$. The x 's are called individual-symbols.
- (2) A finite set of function-symbols, $f^1, \dots, f^n$ ($n \geq 0$). With each function-symbol is given an integer ≥ 0 , the number of its arguments.

¹Kurt Gödel, Über formal unentscheidbare Sätze der Principia Mathematica und verwandte Systeme I, Monatshefte für Mathematik und Physik, vol. 38 (1931), p. 173.

²Alonzo Church, A note on the Entscheidungsproblem, Journal of Symbolic Logic, vol. 1 (1936), p. 40; Correction to A note on the Entscheidungsproblem, ibid., p. 101.

³See, for example, Hilbert's remarks in the Einführung to Grundlagen der Mathematik.

(3) The five symbols $\vee$, $\wedge$, $+$, $\times$, and $'$, called respectively universal, null, plus, times, and negative.

(4) The symbols $\sum$ and $\prod$, called respectively sum and product, and called collectively quantifiers.

Combinations of symbols of the form $f^1(x_{k_1}, x_{k_2}, \dots, x_{k_j})$ ($j \geq 0$) will be called functional-value symbols. Where convenient to avoid a multiplicity of parentheses, the arguments will be written as subscripts.

Not all combinations of the foregoing symbols are admitted as expressions for the purpose of this chapter. First, the individual-symbols must all appear as summed- or multiplied out (that is, none of them enters as a free variable), and none of them is associated with two quantifiers which are distinct as parts of the expression. Second, the function-symbols all appear as free variables. Third, the expression must be capable of evaluation as described in the next paragraph. Such an expression is exemplified by

$$\sum_{x_1} \prod_{x_2} (f_{x_1 x_2}^1 + \prod_{x_3} f_{x_1 x_3}^2) + \sum_{x_4} \vee f_{x_4}^3.$$

The class of all such expressions will be denoted by $\mathfrak{E} = [e]$.

The two-element Boolean algebra $\mathcal{B}^2 = [\odot, \oslash]$ will be used throughout this chapter in evaluating expressions. The evaluation then consists of two parts: first, the choice of a non-null¹ class $\mathfrak{B}$ to be used as the class on which the functions are to be defined and over which the summations and multiplications are to be performed; second, the selection of functions on $\mathfrak{B}$ to $\mathcal{B}^2$ to replace the function-symbols. For each function-symbol

¹The theory of the evaluation of expressions on a null class is trivial, but different from the theory for non-null classes. The capital German $\mathfrak{B}$ will connote non-null class throughout this chapter.

f^1 the function by which it is replaced in this evaluation will be symbolized by

$$(\epsilon(f^1))_{\text{on } \mathcal{B} \dots \text{ to } \mathcal{B}^2}.$$

When the choice of $\mathcal{B}$ and ϵ has been made, interpretation of $\prod$ and $\sum$ respectively as logical product and sum over $\mathcal{B}$ reduces any expression, e , to the value $\bigvee$ or $\bigwedge$; this value will be symbolized by

$$V_{\mathcal{B}\epsilon}(e).$$

Equivalence of expressions may be defined in various ways, whether by reference to the preservation of the value to which an expression reduces or to the class of transformations required to replace one expression by the other. We use in this chapter a definition of the former type: two expressions will be said to be V -equivalent in case they are reduced to the same value by every $(\mathcal{B}, \epsilon)$. The transforms, presently to be defined, of a given expression will all be (by well-known laws) V -equivalent to the given expression, but a definition of equivalence entirely in terms of transformations will not be given.

Certain processes of transformation are valuable for many purposes to simplify a given expression. Thus it is often useful to replace the expression $\prod_{x_1} \prod_{x_2} (f^1_{x_1} + f^2_{x_2})$ by the V -equivalent expression

$$\prod_{x_1} f^1_{x_1} + \prod_{x_2} f^2_{x_2},$$

or such an expression as

$$\prod_{x_1} \sum_{x_2} \prod_{x_3} f^1_{x_3}$$

by the V -equivalent expression $\prod_{x_3} f^1_{x_3}$. These and other types of simplification are very helpful when it is desired actually to

evaluate a given expression, but for the purpose of the existence proofs of this chapter they are not needed, and, in fact, would necessitate the separation of cases which can otherwise be treated together.

This is accomplished by preliminary reduction of the given expression to a normal form characterized as follows: The first quantifier is a Π which applies to all the rest of the expression. The operand of the initial sequence of Π 's is a finite sum of (one or more) terms none of which begins with a Π , each of which has all its quantifiers gathered together at the start, and at least one of which contains a sequence of (one or more) Σ 's followed by a sequence of Π 's. The reduction of any given expression to a V-equivalent normal form can be carried out by elementary processes which are well known.

In the ensuing discussion it is assumed that an expression e is given, and a V-equivalent expression t_0 in normal form has been constructed in any way. It is next desired to describe a process of transformation by which any normal expression is replaced by another normal expression, for the purpose of generating an infinite sequence of transforms $t_1, t_2, \dots$ from t_0 . For this purpose it is helpful to introduce notations for certain parts of a normal expression.

The class of individual-symbols appearing as indices of Π 's of the initial sequence will be denoted by $\mathcal{K}_0$. The summands which follow this sequence of Π 's will be denoted by s^1, s^2 , etc. For each of these terms, s^1 , the classes of individual-symbols which appear as indices of the Σ 's and Π 's of the sequences of like quantifiers in s^1 will be designated with the superscript 1, and distinguished from one another by subscripts showing the order of the sequence of indices in s^1 . Where a sequence t_0 ,

$t_1, \dots$ of normal expressions is under consideration, the subscript j will be added, so that s_j^1 is the i -th term in t_j and $\mathcal{K}_{kj}^1$ is the class of x 's appearing as indices of the k -th sequence of like quantifiers in s_j^1 . Then odd first subscripts of $\mathcal{K}$'s refer to sequences of Σ 's, even to Π 's.

The transform t_{j+1} of any normal expression t_j is constructed in the following manner: Consider any term s_j^1 of t_j such that the symbol Σ appears in s_j^1 . Every function η on $\mathcal{K}_{1j}^1$ to $\mathcal{K}_{0j}$ determines a new term to be obtained from s_j^1 by dropping the initial sequence of Σ 's in s_j^1 and replacing, in the remainder of s_j^1 , each member, x , of $\mathcal{K}_{1j}^1$ by $\eta(x)$. Forming all such terms, and choosing new distinct x 's (from the infinite class $\mathcal{K}$) to replace all individual-symbols in them except members of $\mathcal{K}_{0j}$, we add all such terms obtainable from all terms of t_j to the old terms s_j^1 , and transfer all their initial Π 's to the initial sequence of Π 's; the expression thus obtained is the required t_{j+1} .

An example will illustrate this process. The formula

$$\Pi_{x_1} \Sigma_{x_2} \Pi_{x_3} f_{x_1 x_2 x_3}^1$$

is in normal form. It has only one term,

$$s^1 = \Sigma_{x_2} \Pi_{x_3} f_{x_1 x_2 x_3}^1,$$

following the initial sequence of Π 's. The class $\mathcal{K}_0$ consists of x_1 in this case, and the class $\mathcal{K}_1^1$ of x_2 . Dropping the initial Σ in s^1 , replacing x_2 by x_1 , and choosing the new x_4 to replace x_3 , we obtain the term

$$\Pi_{x_4} f_{x_1 x_1 x_4}^1,$$

and in adjoining this term to s^1 we transfer the $\prod_{x_4}$ to the initial sequence of $\prod$'s; the transform is then

$$\prod_{x_1} \prod_{x_4} (\sum_{x_2} \prod_{x_3} f^1_{x_1 x_2 x_3} + f^1_{x_1 x_1 x_4}).$$

This process of transformation consists primarily in adjoining terms of a sum to the sum itself, and therefore preserves V-equivalence:

(13.1) Theorem. The transforms t_j ($j = 0, 1, \dots$) are all in normal form and all V-equivalent to the given expression e . The classes $\mathcal{K}_{0j}$ and $[all\ s^1_j]$ are properly monotonic-increasing in j .

The sum as to j of all the $\mathcal{K}_{0j}$ will be denoted by $\mathcal{Y}$; this infinite class is enumerated in the process of constructing the transforms, the members of each $\mathcal{K}_{0j}$ being all counted before the first of the remaining x 's of $\mathcal{Y}$.

We next describe how a reduced transform, r_j , is to be constructed from each transform t_j . In every term s^1 of a normal expression the part of the term which follows the last of its quantifiers will be called its matrix; it is an expression, in the sense of Chapter II, in which functional-value symbols serve as letters. Transforming each of these matrices to polynomial form, we define r_j as a sum for all s^1 of all such monomial terms in the matrices as involve no x 's other than members of $\mathcal{K}_{0j}$, provided there are such terms--otherwise $r_j =_{\mathcal{D}} \Lambda$. (In actual calculations it is helpful to use the simplified form for the matrices, but the present discussion applies to any polynomial representations.)

The class of terms of the r_j is evidently monotonic-increasing¹ in j .

Since the reduced transforms are finite Boolean expressions in the sense of Chapter II, the theory there given applies to them. The class of functional-value symbols whose arguments are all members of $\mathcal{U}$ will be denoted by $\mathcal{L}$; the members of this class serve as letters in the r_j . The evaluation of reduced transforms will be confined to the two-element Boolean algebra $\mathcal{B}^2 = [\mathbb{V}, \mathbb{A}]$, and the letter ξ will be reserved in this chapter for a function on $\mathcal{L}$ to $\mathcal{B}^2$ or the extension of such a function, as defined in Chapter II; it assigns the value $\mathbb{V}$ or $\mathbb{A}$ to every r_j .

If a class $\mathcal{B}$ used in evaluating an expression has members in common with the class $\mathcal{U}$ of individual-symbols appearing in the reduced transforms, the evaluations ϵ and ξ will be said to agree in case they yield the same value for every f^1 and every set of arguments x belonging to ${}^{\wedge}(\mathcal{U}, \mathcal{B})$.

Since each reduced transform is obtained from the corresponding transform by dropping some of its terms, it is clear that if a reduced transform is equivalent to $\mathbb{V}$, the corresponding transform (and therefore also the given expression) must assume the value $\mathbb{V}$ in any evaluation:

(13.2) Theorem.

$$j : r_j \sim \mathbb{V} \cdot \epsilon \cdot v_{\mathcal{B}\epsilon}(t_j) = \mathbb{V}$$

The class $\mathcal{B}$ is unrestricted (except that it must be non-null) in this theorem. The converse, that there exists an ϵ which reduces the given expression and its transforms to $\mathbb{A}$ if

¹But not necessarily properly monotonic-increasing; i.e., the class is non-decreasing.

no r_j is equivalent to V , is not true for an unrestricted class, but holds if $\mathcal{B}$ is infinite, as will appear in the next three sections.

14. Existence of an evaluation in which all reduced transforms take the value $\textcircled{V}$. The theorem to be proved in this section is the first part of the converse just mentioned: If no reduced transform is equivalent to V , there exists an evaluation ξ which yields the value $\textcircled{V}$ for every reduced transform.

(14.1) Theorem.

$$r_j \not\sim \wedge (j) : \exists \xi \ni j \cdot \xi(r_j) = \textcircled{V}$$

A geometrical representation of the evaluation functions ξ will be useful in this proof. For every ξ a real number u_1 is to be assigned to each function-symbol f^1 such that $0 \leq u_1 \leq 1$, so that ξ determines a point in or on the boundary of an n -dimensional cube or n -cell of edge unity in a Euclidean n -space, n being the number of the function-symbols.

These real numbers are assigned in the following way. First, for function-symbols f^1 with no arguments, we take $u_1 = 1$ if $\xi(f^1) = \textcircled{V}$, and $u_1 = 0$ if $\xi(f^1) = \textcircled{V}$. Next, for each function-symbol f^1 whose number of arguments is $n_1 > 0$, we denumerate in any way the functional-value symbols based on f^1 and involving only x 's which are members of $\mathcal{U}$. The real number u_1 assigned to f^1 is then the value of the binary radix-fraction¹ whose j -th component is 1 or 0 according as the ξ -function of the j -th functional-value symbol in this denumeration is $\textcircled{V}$ or $\textcircled{V}$.

Every ξ thus determines a unique point $(u_1, u_2, \dots, u_n)$ in the n -cell, which will be said to represent ξ , but a point

¹That is, the binary analogue of a decimal; a Dualbruch.

may represent more functions ξ than one, because certain real numbers admit a double radix-fraction development. But the number of ξ 's represented by any point is finite, being in fact not more than 2^n .

Every condition on ξ determines a set of points in the n -cell. Thus for a function-symbol f^1 without arguments the assignment of a value, $\textcircled{V}$ or $\textcircled{A}$, to $\xi(f^1)$ confines the points representing ξ to one face of the cell. For a function-symbol f^1 with one or more arguments the assignment of a value to one of its functional-value symbols confines u_1 to a finite non-null set of closed intervals, and imposes no further restriction; the point representing ξ is then confined to the corresponding closed part of the n -cell.

For any monomial term m of one of the reduced transforms, the condition $\xi(m) = \textcircled{A}$ is satisfied if and only if the ξ -function of some factor of m is $\textcircled{A}$. The set of points representing [all $\xi \ni \xi(m) = \textcircled{A}$] is then the sum of the sets for the factors of m , and is therefore closed. Similarly the set $\mathcal{G}_j$ of points representing [all $\xi \ni \xi(r_j) = \textcircled{A}$] is closed.

Since every $r_j \neq V$ by hypothesis, it follows by the dual of Theorem (8.2) that all the $\mathcal{G}_j$ are non-null. These sets are evidently bounded, and they are also monotonic-decreasing in j because each r_j has at least all the terms of its predecessors. The intersection $\mathcal{G}$ of all $\mathcal{G}_j$ is therefore non-null and closed.

For any point of $\mathcal{G}$ the set of functions ξ which are represented by the point is non-null and finite, and for every j the subset of these functions for which $\xi(r_j) = \textcircled{A}$ is non-null; these sets are monotonic-decreasing in j . Some function ξ belongs, therefore, to all these sets, and serves as the ξ of the theorem.

15. Existence of an evaluation reducing the initial transform to $\mathbb{A}$. The theorem of this section is the second part of the converse mentioned in Section 13. If an evaluation ϵ for the class $\mathcal{Y}$ agrees with an evaluation ξ for which every reduced transform takes the value $\mathbb{A}$, then ϵ reduces the initial transform to $\mathbb{A}$.

(15.1) Theorem.

$$\begin{aligned} \mathcal{B} &= \mathcal{Y} \cdot \xi(r_j) = \mathbb{A} \quad (j) \cdot \epsilon \text{ agrees with } \xi \cdot \\ v_{3\epsilon}(t_0) &= \mathbb{A} \end{aligned}$$

In an evaluation to the 2-element Boolean algebra $[\mathbb{A}, \mathbb{A}]$, a product takes the value $\mathbb{A}$ if and only if some factor takes this value. Therefore, since t_0 begins with a sequence of Π 's, it is to be proved that the x 's of the class $\mathcal{X}_{00}$ (that is the x 's associated with this sequence of Π 's) can be replaced, wherever they occur in the part of t_0 that follows the initial Π 's, by members of $\mathcal{B}$ in such a way that the expression reduces to $\mathbb{A}$. It will be proved that the identity-function $\eta^0(x) = x$ ($x \in \mathcal{X}_{00}$) serves this purpose. This substitution leaves the terms s_0^1 of t_0 unchanged.

Since a sum takes the value $\mathbb{A}$ if and only if every term takes this value, it is to be proved that every term s_0^1 reduces to $\mathbb{A}$, that is, that for every η^1 on $\mathcal{X}_{10}^1$ to $\mathcal{B}$ there exists an η^2 on $\mathcal{X}_{20}^1$ to $\mathcal{B}$ such that the substitution of $\eta^1(x)$ for every x of $\mathcal{X}_{10}^1$ and of $\eta^2(x)$ for every x of $\mathcal{X}_{20}^1$ reduces the part of s_0^1 which follows its first sequence of Π 's to an expression which on further substitution in a similar manner of members of $\mathcal{B}$ for members of $\mathcal{X}_{30}^1, \mathcal{X}_{40}^1$, etc., ultimately reduces the term to $\mathbb{A}$.

The function η^{2k} on $\mathcal{X}_{2k,0}^1$ to $\mathcal{B}$ ($k = 1, 2, \dots$) is to be chosen in the following way in terms of the preceding η 's.

Let t_j be the first transform for which all $\eta(x)$

$(x(\mathcal{X}_{10}^1, \mathcal{X}_{20}^1, \dots, \mathcal{X}_{2k-1,0}^1))$ are members of $\mathcal{X}_{0j}$. To proceed by induction on k it will be supposed that t_j has a term $s_j^{m_1}$ derived ultimately from s_0^1 in such a way that $\mathcal{X}_{1j}^{m_1}$ corresponds to $\mathcal{X}_{2k-1,0}^1$, $\mathcal{X}_{2j}^{m_1}$ to $\mathcal{X}_{2k,0}^1$, etc., and every x of

$$(\mathcal{X}_{00}, \mathcal{X}_{10}^1, \mathcal{X}_{20}^1, \dots, \mathcal{X}_{2k-2,0}^1)$$

appearing in the matrix of s_0^1 is replaced in $s_j^{m_1}$ by $\eta(x)$. Let ζ denote the inverse of the product of the replacements of old x 's by new by which $\mathcal{X}_{1j}^{m_1}$ was obtained from $\mathcal{X}_{2k-1,0}^1$ in the process of constructing the transforms. The values of the function

$$(\eta^{2k-1}(\zeta(x)) | x \mathcal{X}_{1j}^{m_1})$$

are all members of $\mathcal{X}_{0j}$, so that this function determines a term $s_{j+1}^{m_2}$ of t_{j+1} . The x 's adopted in the construction of $s_{j+1}^{m_2}$ to replace the members of $\mathcal{X}_{2j}^{m_1}$ are the respective values to be chosen for the function η^{2k} .

Evidently the term $s_{j+1}^{m_2}$ has the property with which the induction started, and, since all terms of t_0 have this property by the choice of η^0 , there is ultimately a transform t_h containing the term obtained from the matrix of s_0^1 by the substitution of $\eta(x)$ for each x . Since this term involves no x 's except members of $\mathcal{Y}$, it is a term of the reduced transform r_h , and is therefore by hypothesis reduced to $\mathcal{O}$ by the evaluation ξ , and therefore also by ϵ , which agrees with ξ .

16. The equivalence theorem. The foregoing results are here combined in a theorem which provides conditions under which an expression can be evaluated to $\mathbb{A}$: For any expression e and its sequence of reduced transforms r_j , the expression can be evaluated to $\mathbb{A}$ on an infinite class if and only if none of the r_j is equivalent to V .

(16.1) Theorem.

$$\begin{aligned} & (1) \quad \mathcal{B}^{-II} \cdot). \quad \exists \epsilon \ni v_{\mathcal{B}\epsilon}(e) = \mathbb{A} \\ \equiv & (2) \quad \exists(\mathcal{B}^{-II}, \epsilon) \ni v_{\mathcal{B}\epsilon}(e) = \mathbb{A} \\ \equiv & (3) \quad \exists(\mathcal{B}, \epsilon) \ni v_{\mathcal{B}\epsilon}(e) = \mathbb{A} \\ \equiv & (4) \quad r_j \not\sim V(j) \end{aligned}$$

It suffices to prove (1) $\cdot$) $\cdot$ (2) $\cdot$) $\cdot$ (3) $\cdot$) $\cdot$ (4) $\cdot$) $\cdot$ (1). The first two of these implications are obvious, and the third is a restatement of Theorem (13.2) since e is V -equivalent to t_j . To prove that (1) follows from (4), let $\mathcal{B}_0$ be a denumerably infinite subclass of $\mathcal{B}$ and adopt a one-to-one reciprocal correspondence between $\mathcal{B}_0$ and $\mathcal{Y}$. Then by Theorem (15.1) any ϵ_0 which agrees with the ξ of Theorem (14.1) yields the result $v_{\mathcal{B}_0\epsilon_0}(e) = \mathbb{A}$. To extend this result to $\mathcal{B}$ it is sufficient to use any η on $\mathcal{B}$ to $\mathcal{B}_0$ such that $\eta(z_0) = z_0$ for every member of $\mathcal{B}_0$, and define $\epsilon(f^1)$ for any arguments z as the value of $\epsilon_0(f^1)$ for $\eta(z)$.

If the class $\mathcal{B}$ in condition (1) is taken to be denumerably infinite, the implication (3) $\cdot$) $\cdot$ (1) yields Löwenheim's theorem¹ that a formula of the Prädikatenkalkül can be reduced to $\mathbb{A}$ on a denumerable class if at all.

¹Leopold Löwenheim, "Über Möglichkeiten im Relativkalkül, Mathematische Annalen, vol. 76 (1915), p. 447.

The fact that condition (4) is equivalent to the others is somewhat analogous to Skolem's theorem¹ on the completeness of the Prädikatenkalkül, but is stronger² in that it refers to a fixed process of transformation rather than to the totality of laws of the calculus.

17. Example. The method will be illustrated on a simple problem, namely the formula $\bar{\mathfrak{F}}$ discussed by Hilbert and Bernays on page 123 of volume 1 of Grundlagen der Mathematik. Changing the notation to that of this paper and reducing the expression to normal form we obtain

$$\begin{aligned}
 (17.1) \quad t_0 = & \prod_{x_1} (\sum_{x_2} f^1(x_2, x_2) \\
 & + \sum_{x_3} \sum_{x_4} \sum_{x_5} f^1(x_3, x_4) f^1(x_4, x_5) f^1(x_3, x_5) \\
 & + \sum_{x_6} \prod_{x_7} f^1(x_6, x_7)).
 \end{aligned}$$

It will be proved that no reduced transform r_j of this expression is equivalent to V .

First, by the construction of the transforms t_j , every term of r_j which is derived from the third term of t_0 has the property that the first of its arguments precedes the second in the denumeration of $\mathfrak{N}$.

Second, all terms of r_j derivable from the second term of t_0 have the following property: there exist two individual-

¹Thoralf Skolem, Logisch-Kombinatorische Untersuchungen über die Erfüllbarkeit oder Beweisbarkeit Mathematische Sätze nebst einem Theoreme über dichte Mengen, Skrifter utgitt av Videnskapsselskapet i Kristiania (Oslo), I. Matematisk-Naturvidenskabelig Klasse (1920), No. 4, p. 6.

²The result of Skolem has been strengthened in another direction by Gödel, namely in the scope of the formulas dealt with: Kurt Gödel, Die Vollständigkeit der Axiome des logischen Funktionenkalküls, Monatshefte für Mathematik und Physik, vol. 37 (1930), p. 349.

symbols, x_{1_1} and x_{1_2} , such that $f^{1'}(x_{1_1}, x_{1_2})$ appears as a factor of the term, and there appears (or can be made to appear by a rearrangement of the factors) the product of a sequence of functional-value symbols,

$$f^1(x_{1_1}, x_{k_1}), f^1(x_{k_1}, x_{k_2}), \dots, f^1(x_{k_m}, x_{1_2})$$

whose arguments begin with x_{1_1} and end with x_{1_2} , such that the first argument of each factor except the first is the same as the second argument of its predecessor.

Terms of r_j which are derived from the first and third terms of t_0 do not have this property, but they can be admitted by weakening it in two ways: First, the factor $f^{1'}(x_{1_1}, x_{1_2})$ may be omitted if $x_{1_1} = x_{1_2}$. Second, any subsequences $f^1(x_{k_1}, x_{k_2})f^1(x_{k_2}, x_{k_3})\dots f^1(x_{k_{h-1}}, x_{k_h})$ of consecutive factors may be omitted for which x_{k_1} precedes x_{k_h} in the denumeration of $\mathcal{G}$.

By Theorem (10.3), if $r_j \sim V$, it can be reduced to a polynomial one of whose terms is V by the adjunction of σ -terms obtained by combining pairs of terms in accordance with Theorem (10.2). By part (1) of Definition (10.1), it is seen that if two monomials have the weakened property described above, and if they admit such a σ -term, this term also has the weakened property. Finally, V has not this property, for if so, then $x_{1_1} = x_{1_2}$ and x_{1_1} precedes x_{1_2} in the denumeration of $\mathcal{G}$.

VITA

Archie Blake was born on November 24, 1906, in Washington, D. C. After passing through the public schools of Chicago, Illinois, he entered the University of Chicago in October, 1924. He was graduated with the S.B. degree in June, 1929. After teaching physics for a year at the Georgia School of Technology, he returned in the fall of 1930 to the University of Chicago, receiving the S.M. degree in mathematics in June, 1931.

Since then he has been working in mathematical seismology at the U. S. Coast and Geodetic Survey, returning at intervals to the University of Chicago to pursue his studies in pure mathematics.

The writer takes this opportunity to acknowledge his indebtedness to all of his teachers, and to two of them especially: to Professor A. C. Lunn for much inspiration in the study of mathematical physics, and to Professor R. W. Barnard for invaluable help in the preparation both of this dissertation and of the writer's master's thesis. Indirectly he is also indebted to E. H. Moore, whose methods and notations have been of great help to the writer both in reading others' work and in formulating his own.

